



Digitized by the Internet Archive
in 2024

https://archive.org/details/bwb_P9-DSO-984



H. Minkowski

GESAMMELTE ABHANDLUNGEN

VON

HERMANN MINKOWSKI

UNTER MITWIRKUNG VON

ANDREAS SPEISER UND HERMANN WEYL

HERAUSGEGEBEN VON

DAVID HILBERT

CHELSEA PUBLISHING COMPANY

NEW YORK, N. Y.

THE PRESENT WORK IS A REPRINT, TEXTUALLY UNALTERED
EXCEPT FOR THE CORRECTION OF ERRATA, OF THE WORK
GESAMMELTE ABHANDLUNGEN, VON HERMANN MINKOWSKI.
THE TWO VOLUMES THAT CONSTITUTE THE ORIGINAL WORK
ARE NOW REPRINTED IN ONE VOLUME. THE TWO FOLD-OUT
PLATES THAT APPEARED AT THE END OF THE SECOND VOLUME
ARE INCORPORATED INTO THE TEXT, AS EIGHT HALF-TONE
PLATES, AT THE APPROPRIATE PLACES IN VOLUME TWO.

ORIGINALLY PUBLISHED, IN TWO VOLUMES, LEIPZIG, 1911

REPRINTED, TWO VOLUMES IN ONE, NEW YORK, N. Y., 1967

LIBRARY OF CONGRESS CATALOG CARD NO. 66-28570

PRINTED IN THE UNITED STATES OF AMERICA

INHALT DES ERSTEN BANDES.

Gedächtnisrede auf H. Minkowski, von D. Hilbert	Seite V
---	------------

Zur Theorie der quadratischen Formen.

I. Grundlagen für eine Theorie der quadratischen Formen mit ganzzahligen Koeffizienten.	3
(In französischer Sprache unter dem Titel: Mémoire sur la théorie des formes quadratiques, in den Mémoires présentés par divers savants à l'Académie des Sciences de l'Institut national de France, Tome XXIX, No. 2; 1884.)	
II. Sur la réduction des formes quadratiques positives quaternaires	145
(Comptes rendus de l'Académie des Sciences, Paris, t. 96, pp. 1205—1210; 1883.)	
III. Über positive quadratische Formen	149
(Crelles Journal für die reine und angewandte Mathematik, Bd. 99, S. 1—9; 1886.)	
IV. Untersuchungen über quadratische Formen. Bestimmung der Anzahl verschiedener Formen, welche ein gegebenes Genus enthält.	157
(Inauguraldissertation, Königsberg 1885; Acta Mathematica, Bd. 7, S. 201—258; 1885.)	
V. Über den arithmetischen Begriff der Äquivalenz und über die endlichen Gruppen linearer ganzzahliger Substitutionen. . . .	203
(Crelles Journal für die reine und angewandte Mathematik, Bd. 100, 449—458; 1887.)	
VI. Zur Theorie der positiven quadratischen Formen	212
(Crelles Journal für die reine und angewandte Mathematik, Bd. 101, S. 196—202; 1887.)	
VII. Über die Bedingungen, unter welchen zwei quadratische Formen mit rationalen Koeffizienten ineinander rational transformiert werden können (Auszug aus einem von Herrn H. Minkowski in Bonn an Herrn Adolf Hurwitz gerichteten Brief).	219
(Crelles Journal für die reine und angewandte Mathematik, Bd. 106, S. 5—26; 1890.)	
Zur Geometrie der Zahlen.	
VIII. Über die positiven quadratischen Formen und über kettenbruchähnliche Algorithmen.	243
(Crelles Journal für die reine und angewandte Mathematik, Bd. 107, S. 278—297; 1891.)	

IX. Théorèmes arithmétiques (Extrait d'une lettre de M. H. Minkowski à M. Hermite)	261
(Comptes rendus de l'Académie des Sciences, Paris, t. 112, pp. 209—212; 1891.)	
X. Über Geometrie der Zahlen (Bericht über einen Vortrag zu Halle) .	264
(Verhandlungen der 64. Naturforscher- und Ärzteversammlung zu Halle, 1891, S. 13, und Jahresbericht der Deutschen Mathematiker-Vereinigung, Bd. 1, S. 64—65; 1892.)	
XI. Extrait d'une lettre adressée à M. Hermite.	266
(Bulletin des Sciences mathématiques, 2 ^e série, t. XVII, pp. 24—29; 1893.)	
XII. Über Eigenschaften von ganzen Zahlen, die durch räumliche Anschauung erschlossen sind	271
(Mathematical Papers read at the international Mathematical Congress held in connection with the world's Columbian Exposition Chicago, 1893, pp. 201—207; ferner unter dem Titel: Sur les propriétés des nombres entiers qui sont dérivées de l'intuition de l'espace, von L. Laugel ins Französische übersetzt, in Nouvelles Annales de Mathématiques, 3 ^e série, t. XV, pp. 393—403; 1896.)	
XIII. Zur Theorie der Kettenbrüche	278
(Von L. Laugel ins Französische übersetzt unter dem Titel: Généralisation de la théorie des fractions continues, in Annales de l'École Normale supérieure, 3 ^e série, t. XIII, pp. 41—60; 1896.)	
XIV. Ein Kriterium für die algebraischen Zahlen	293
(Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen, mathematisch-physikalische Klasse, 1899, S. 64—88.)	
XV. Zur Theorie der Einheiten in den algebraischen Zahlkörpern	316
(Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen, mathematisch-physikalische Klasse, 1900, S. 90—93.)	
XVI. Über die Annäherung an eine reelle GröÙe durch rationale Zahlen	320
(Mathematische Annalen, Bd. 54, S. 91—124; 1901.)	
XVII. Quelques nouveaux théorèmes sur l'approximation des quantités à l'aide de nombres rationnels	353
(Bulletin des Sciences mathématiques, 2 ^e série, t. XXV, pp. 72—76; 1901.)	
XVIII. Über periodische Approximationen algebraischer Zahlen. . .	357
(Acta Mathematica, Bd. 26, S. 333—351; 1902.)	

Hermann Minkowski.

Gedächtnisrede, gehalten in der öffentlichen Sitzung der K. Gesellschaft der
Wissenschaften zu Göttingen am 1. Mai 1909

von

David Hilbert.

(Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen. 1909.)

Einen schweren unermeßlichen Verlust haben zu Beginn des Jahres 1909 unsere Gesellschaft, unsere Universität, die Wissenschaft und wir alle persönlich erlitten: durch ein hartes Geschick wurde uns jäh entrissen unser Kollege und Freund Hermann Minkowski im Vollbesitz seiner Lebenskraft, aus der Mitte freudigsten Wirkens, von der Höhe seines wissenschaftlichen Schaffens.

Seinem Andenken widmen wir diese Stunde.

Hermann Minkowski wurde am 22. Juni 1864 zu Alexoten in Rußland geboren, kam als Knabe nach Deutschland und trat Oktober 1872 im Alter von $8\frac{1}{4}$ Jahren in die Septima des Altstädtischen Gymnasiums zu Königsberg i. Pr. ein. Da er von sehr rascher Auffassung war und ein vortreffliches Gedächtnis hatte, wurde er auf mehreren Klassen in kürzerer als der vorgeschriebenen Zeit versetzt und verließ das Gymnasium schon März 1880 — noch als Fünfzehnjähriger — mit dem Zeugnis der Reife.

Ostern 1880 begann Minkowski seine Universitätsstudien. Insgesamt hat er 5 Semester in Königsberg, vornehmlich bei Weber und Voigt, und 3 Semester in Berlin studiert, wo er die Vorlesungen von Kummer, Kronecker, Weierstraß, Helmholtz und Kirchhoff hörte.

Seine Befähigung zur Mathematik zeigte sich früh; fiel ihm doch im ersten Semester bereits für die Lösung einer mathematischen Aufgabe eine Geldprämie zu, auf die er freilich zugunsten eines armen Mitschülers verzichtete, so daß sein frühzeitiger Erfolg zu Hause gar nicht bekannt wurde — eine kleine Begebenheit, die zugleich die Bescheidenheit und Herzensgüte kennzeichnet, wie er sie sein ganzes Leben hindurch allen Menschen gegenüber, die ihm näher kamen, betätigt hat.

Sehr bald begann Minkowski tiefgehende und gründliche mathematische Studien. Ostern 1881 hatte die Pariser Akademie das Problem der Zerlegung der ganzen Zahlen in eine Summe von fünf Quadraten als Preisthema gestellt. Dieses Thema griff der siebzehnjährige Student mit aller

Energie an und löste die gestellte Aufgabe aufs glänzendste, indem er weit über das Preisthema hinaus die allgemeine Theorie der quadratischen Formen, insbesondere ihre Einteilung in Ordnungen und Geschlechter — zunächst sogar für beliebigen Trägheitsindex — entwickelte*). Es ist erstaunlich, welch sichere Herrschaft Minkowski schon damals über die algebraischen Methoden, insbesondere die Elementarteilertheorie, sowie über die transzendenten Hilfsmittel wie die Dirichletschen Reihen und die Gaußschen Summen besaß, — Kenntnisse, die noch heute lange nicht allgemeines Eigentum der Mathematiker geworden sind, die aber freilich zur erfolgreichen Inangriffnahme des Pariser Preisthemas eine notwendige Voraussetzung bildeten. Hören wir, wie Minkowski selbst in dem Begleitschreiben zu seiner der Pariser Akademie eingereichten Arbeit sich ausspricht**): „Durch die von der Académie des Sciences gestellte Aufgabe angeregt“, so schreibt der jugendliche Student, „unternahm ich eine genauere Untersuchung der allgemeinen quadratischen Formen mit ganzzahligen Koeffizienten. Ich ging dabei von dem natürlichen Gedanken aus, daß die Zerlegung einer Zahl in eine Summe von fünf Quadraten in ähnlicher Weise von den quadratischen Formen mit vier Variablen abhängen würde, wie bekanntlich die Zerlegung einer Zahl in eine Summe von drei Quadraten von den quadratischen Formen mit zwei Variablen abhängt. Diese Untersuchung hat mir in der Tat die gewünschten Resultate über die Zerlegung einer Zahl in eine Summe von fünf Quadraten geliefert. Indessen erscheinen diese Resultate bei der großen Allgemeinheit der von mir gefundenen Sätze nicht überall als das eigentliche Hauptziel der vorliegenden Arbeit; sie stellen vielmehr nur ein Beispiel für die gewonnenen umfangreichen Theorien dar. Wenn daher viele der nachfolgenden Betrachtungen nicht immer unmittelbar auf das Thema der Preisfrage hinweisen, so wage ich dennoch zu hoffen, daß die Akademie nicht der Ansicht sein werde, ich würde mehr gegeben haben, wenn ich weniger gegeben hätte.“ Mit dem Motto: „Rien n'est beau que le vrai, le vrai seul est aimable“ reichte der noch nicht Achtzehnjährige am 30. Mai 1882 die Arbeit der Pariser Akademie ein. Obwohl dieselbe, entgegen den Bestimmungen der Akademie, in deutscher Sprache abgefaßt war, so erkannte die Akademie dennoch unter ausdrücklicher Betonung des exzeptionellen Falles auf Zuerteilung des vollen Preises, da — wie es im Kommissionsbericht heißt — eine Arbeit von solcher Bedeutung nicht wegen einer Irregularität der Form von der

*) „Mémoire sur la théorie des formes quadratiques à coefficients entiers.“ Mémoires présentés par divers savants à l'Académie des Sciences de l'Institut national de France, T. XXIX. No. 2 (1884). Unter dem Titel „Grundlagen für eine Theorie der quadratischen Formen mit ganzzahligen Koeffizienten“, diese Ges. Abhandlungen, Bd. I, S. 3—144. **) Vgl. diese Ges. Abhandlungen, Bd. I, S. 4.

Bewerbung auszuschließen sei, und erteilte ihm im April 1883 den Grand Prix des Sciences Mathématiques.

Als die Zuerkennung des Akademiepreises an Minkowski in Paris bekannt wurde, richtete die dortige chauvinistische Presse gegen ihn die unbegründetsten Angriffe und Verdächtigungen. Die französischen Akademiker C. Jordan und J. Bertrand stellten sich sofort rückhaltlos auf die Seite Minkowskis. „Travaillez, je vous prie, à devenir un géomètre éminent.“ In dieser Mahnung des großen französischen Mathematikers C. Jordan an den jungen deutschen Studenten gipfelte die bei diesem Anlaß zwischen C. Jordan und Minkowski geführte Korrespondenz, — eine Mahnung, die Minkowski treulich beherzigt hat; begann doch nun für ihn eine arbeitsfrohe und publikationsreiche Zeit.

Gauß hat in seinen *Disquisitiones arithmeticae* die Theorie der binären quadratischen Formen mit ganzzahligen Koeffizienten und damit zugleich den wesentlichen Inhalt der heutigen Theorie der quadratischen Zahlkörper geschaffen. Nach zwei verschiedenen Richtungen hin war die Verallgemeinerung der Gaußschen Theorie möglich: einmal als Theorie der quadratischen Formen mit beliebig vielen Variablen und dann als Theorie der zerlegbaren Formen höherer Ordnung, d. h. als Theorie der Zahlkörper von beliebigem Grade. Durch das Pariser Preisthema war Minkowski zunächst auf die erstere Verallgemeinerung der Gaußschen Theorie hingewiesen: in der Tat sehen wir Minkowski in den folgenden Jahren ausschließlich seine ganze Arbeitskraft dem Studium der *Theorie der quadratischen Formen* und der aufs engste damit zusammenhängenden Fragen widmen. Die Gaußsche Theorie der quadratischen Formen hatte eine wesentliche Ergänzung durch Dirichlet erfahren, indem es diesem gelungen war, auf Grund einer ihm eigentümlichen transzendenten Methode für die Anzahl der Klassen binärer quadratischer Formen mit gegebener Determinante geschlossene Ausdrücke aufzustellen. Es lag nahe, diese Methode nach jenen beiden oben gekennzeichneten Richtungen hin zu verallgemeinern. Nach letzterer Richtung hin, nämlich für die Theorie der algebraischen Zahlkörper, war jene Verallgemeinerung der Dirichletschen Methode bereits von Kummer und in allgemeinsten Weise von Dedekind vorgenommen worden; in ersterer Richtung aber, nämlich für das Problem der quadratischen Formen von beliebig vielen Variablen, lagen nur einige Vorarbeiten von St. Smith, jenem schon bejahrten englischen Zahlentheoretiker, vor, welcher auch bei der Bewerbung um den Pariser Preis Minkowskis Konkurrent gewesen war. Minkowski führte nun die Bestimmung der Anzahl der in einem Geschlecht enthaltenen Klassen quadratischer Formen von beliebig vielen Variablen — denn darauf spitzt sich das in Frage kommende Problem zu — nach der von Dirichlet für binäre

quadratische Formen angewandten transzendenten Methode durch. Die hierbei gefundenen Resultate bilden den wesentlichen Inhalt der Inaugural-Dissertation*), auf Grund deren Minkowski am 30. Juli 1885 von der philosophischen Fakultät in Königsberg zum Doktor promoviert wurde.

Wie glücklich die Ideen des jugendlichen Minkowski auch auf anderem als rein zahlentheoretischem Gebiete waren, ersehen wir aus der bei dieser Gelegenheit von ihm aufgestellten These, die so lautete: „Es ist nicht wahrscheinlich, daß eine jede positive Form sich als eine Summe von Formenquadraten darstellen läßt.“ Es fiel mir als Opponent die Aufgabe zu, bei der öffentlichen Promotion diese These anzugreifen. Die Disputation schloß mit meiner Erklärung, ich sei durch seine Ausführungen überzeugt, daß es wohl schon im ternären Gebiete solch merkwürdige Formen geben möchte, die so eigensinnig seien, positiv zu bleiben, ohne sich doch eine Darstellung als Summe von Formenquadraten gefallen zu lassen. Die Minkowskische These war für mich später die Veranlassung, die Untersuchung der Frage aufzunehmen und für die in der These ausgesprochene Vermutung den strengen Nachweis zu erbringen. Es stellte sich außerdem späterhin heraus, daß das Problem der Darstellung definiter Formen durch Formenquadrate auch bei der Frage nach der Möglichkeit geometrischer Konstruktionen mittels gewisser elementarer Hilfsmittel eine interessante Rolle spielt und andererseits mit gewissen tieferen Problemen über die Darstellbarkeit algebraischer Zahlen als Summen von Quadraten zusammenhängt. Auch von anderer Seite ist seitdem das Problem aufgenommen worden und hat zu interessanten speziellen Ergebnissen geführt.

Angeregt durch eine von Kronecker gestellte Forderung, die eine schärfere Fassung des arithmetischen Begriffs der Äquivalenz von Formen betraf, gelangte Minkowski zu der interessanten Frage nach dem Verhalten linearer ganzzahliger Substitutionen von beliebiger Variablenzahl im Sinne der Kongruenz nach einem beliebigen Modul**). Minkowski gewann dabei den anwendungsreichen Satz, daß eine homogene lineare ganzzahlige Substitution mit n Variablen von einer endlichen Ordnung, die nach einem ganzzahligen Modul ≥ 3 der identischen Substitution kongruent ausfällt, selbst notwendig die identische Substitution ist. Mit Hilfe dieses Satzes

*) Untersuchungen über quadratische Formen. I. Bestimmung der Anzahl verschiedener Formen, welche ein gegebenes Genus enthält. *Acta Mathematica*, Bd. 7 (1885), S. 201—258. Diese Ges. Abhandlungen, Bd. I, S. 157—202.

**) Ueber den arithmetischen Begriff der Aequivalenz und über die endlichen Gruppen linearer ganzzahliger Substitutionen. *Crelles Journal*, Bd. 100 (1887), S. 449—458. Diese Ges. Abhandlungen, Bd. I, S. 203—211. Zur Theorie der positiven quadratischen Formen. *Crelles Journal*, Bd. 101 (1887), S. 196—202. Diese Ges. Abhandlungen, Bd. I, S. 212—218.

gelingt es Minkowski unter anderem zu zeigen, daß die Ordnung jeder endlichen Gruppe von homogenen linearen ganzzahligen Substitutionen mit n Variablen stets ein Divisor der Zahl

$$2^n(2^n - 1)(2^n - 2) \dots (2^n - 2^{n-1})$$

ist, und desgleichen stellt er eine nur von n abhängige Zahl auf, in welcher notwendig allemal die Anzahl der ganzzahligen Substitutionen aufgehen muß, die eine definite quadratische Form mit n Variablen in sich selbst überführen. Die beiden Abhandlungen, welche diese Resultate entwickeln, reichte er der philosophischen Fakultät in Bonn als Habilitationsschrift ein; April 1887 erteilte ihm diese die *venia legendi* für Mathematik.

Noch eine Arbeit Minkowskis sei hier genannt, die ich der Jugend-epoche seines mathematischen Schaffens zuzähle, da sie ebenfalls ausschließlich das Gebiet der quadratischen Formen betrifft; es ist diejenige*), in welcher Minkowski die Bedingungen dafür aufstellt, daß eine quadratische Form mit rationalen Zahlenkoeffizienten sich vermöge einer linearen Substitution mit rationalen Zahlenkoeffizienten in eine andere ebensolche quadratische Form oder in ein rationales Vielfaches einer solchen Form transformieren läßt. Als äußerer Anlaß dazu diente ihm eine von Hurwitz und mir gemeinsam verfaßte Arbeit über ternäre diophantische Gleichungen vom Geschlechte Null. Die Untersuchung von Hurwitz und mir hatte ergeben, daß jede ternäre diophantische Gleichung vom Geschlechte Null durch eine rationale eindeutig umkehrbare Transformation in eine quadratische Gleichung übergeführt werden kann; die weiter entstehenden Fragen, insbesondere die Frage nach den Kriterien dafür, daß eine quadratische diophantische Gleichung bei beliebiger Variablenzahl durch rationale Zahlen lösbar ist, finden durch Minkowski ihre vollständige Erledigung; doch gestaltet sich noch darüber hinaus die Bearbeitung des Problems durch Minkowski zu einer vollständigen Invariantentheorie der quadratischen Formen im zahlentheoretischen Sinne.

Nunmehr beginnt für Minkowskis mathematische Produktion die reichste und bedeutendste Epoche; seine bisher auf das spezielle Gebiet der quadratischen Formen gerichteten Untersuchungen erhalten mehr und mehr den großen Zug ins Allgemeine und gipfeln schließlich in der Schaffung und dem Ausbau der Lehre, für die er selbst den treffenden Namen „*Geometrie der Zahlen*“ geprägt hat und die er in dem großartig angelegten Werke gleichen Titels dargestellt hat.

Das Problem, aus den unendlich vielen Formen einer Klasse durch

*) Ueber die Bedingungen, unter welchen zwei quadratische Formen mit rationalen Coefficienten in einander rational transformiert werden können. Crelles Journal, Bd. 106 (1890), S. 5—26. Diese Ges. Abhandlungen, Bd. I, S. 219—239.

bestimmte Ungleichheitsbedingungen eine einzige auszusondern, d. h. das Problem der Reduktion der quadratischen Formen, hatte Minkowski schon wiederholt beschäftigt. Vor allem ergriffen ihn die berühmten Briefe, die 1850 Ch. Hermite über diesen Gegenstand an Jacobi gerichtet hatte, und insbesondere der dort von Hermite aufgestellte Satz, daß die kleinste von Null verschiedene Größe, die durch eine positive quadratische Form von n Variablen mit der Determinante 1 mittels ganzer Zahlen darstellbar ist, niemals einen gewissen, nur von der Zahl n abhängigen Betrag übersteigt. Durch die Beschäftigung mit diesem Satze wurde Minkowski zu Betrachtungen veranlaßt, auf die wir ein wenig näher eingehen müssen.

Wir denken uns nach Minkowski dasjenige würfelförmig angeordnete, den ganzen Raum erfüllende Punktsystem, welches entsteht, wenn man den rechtwinkligen Koordinaten x, y, z alle ganzzahligen Werte erteilt. Minkowski nannte ein solches Punktsystem ein Zahlengitter. Bedeutet nun $F(x, y, z)$ eine homogene positive quadratische Form von x, y, z mit der Determinante 1, so stellt die Gleichung $F(x, y, z) = c$ für irgendeinen positiven Wert der Konstanten c ein bestimmtes Ellipsoid mit dem Nullpunkt als Mittelpunkt dar. Wir denken uns nun um jeden Punkt des Zahlengitters als Mittelpunkt ein diesem Ellipsoid kongruentes und ähnlich gelegenes Ellipsoid konstruiert: ist dann der Wert der Konstanten c genügend klein, so werden diese Ellipsoide offenbar sämtlich völlig voneinander getrennt liegen. Der größte Wert von c , bei welchem dies noch der Fall ist und die Ellipsoide demnach einander nur in einzelnen Punkten berühren, sei $\frac{1}{4}M$. Da bei dieser Raumerfüllung auf je einen Würfel mit der Kantenlänge 1 je eines der Ellipsoide kommt, so folgt leicht, daß der Inhalt des Ellipsoides $F(x, y, z) = \frac{1}{4}M$ notwendig kleiner als der Inhalt jenes Würfels ausfällt, d. h. es ist gewiß

$$\frac{4\pi}{3} \sqrt{\left(\frac{M}{4}\right)^3} < 1.$$

Andererseits ist leicht zu erkennen, daß das Ellipsoid $F(x, y, z) = M$ gewiß außer dem Nullpunkt keinen Punkt des Zahlengitters in seinem Innern enthält; liegen doch auf seiner Oberfläche gerade noch diejenigen Gitterpunkte, die die Mittelpunkte der das Ellipsoid $F(x, y, z) = \frac{1}{4}M$ berührenden Ellipsoide sind, d. h. M ist der kleinste von Null verschiedene, durch ganze Zahlen darstellbare Wert der quadratischen Form, und jene Ungleichung liefert für dieses Minimum die obere Schranke

$$M < \sqrt[3]{\frac{6^2}{\pi^2}}.$$

Dieser Beweis eines tiefliegenden zahlentheoretischen Satzes ohne rechnerische Hilfsmittel wesentlich auf Grund einer geometrisch anschau-

lichen Betrachtung ist eine Perle Minkowskischer Erfindungskunst. Bei der Verallgemeinerung auf Formen mit n Variablen führt der Minkowskische Beweis auf eine natürlichere und weit kleinere obere Schranke für jenes Minimum M , als sie bis dahin Hermite gefunden hatte. Noch wichtiger aber als dies war es, daß der wesentliche Gedanke des Minkowskischen Schlußverfahrens nur die Eigenschaft des Ellipsoides, daß dasselbe eine konvexe Figur ist und einen Mittelpunkt besitzt, benutzte und daher auf beliebige konvexe Figuren mit Mittelpunkt übertragen werden konnte. Dieser Umstand führte Minkowski zum ersten Male zu der Erkenntnis, daß überhaupt der *Begriff des konvexen Körpers* ein fundamentaler Begriff in unserer Wissenschaft ist und zu deren fruchtbarsten Forschungsmitteln gehört.

Ein konvexer (nirgends konkaver) Körper ist nach Minkowski als ein solcher Körper definiert, der die Eigenschaft hat, daß, wenn man zwei seiner Punkte ins Auge faßt, auch die ganze geradlinige Strecke zwischen denselben zu dem Körper gehört.

Die Bedeutung des Begriffs des konvexen Körpers für die Grundlagen der Geometrie beruht in dem engen Zusammenhange, der, wie Minkowski erkannte, zwischen diesem Begriff und dem fundamentalen Satze Euklids besteht, wonach im Dreiecke die Summe zweier Seiten stets größer als die dritte Seite ist. Dieser Satz Euklids, welcher ja lediglich von elementaren, aus den Axiomen unmittelbar entnommenen Begriffen handelt, folgt bei Euklid aus dem Axiom von der Kongruenz zweier Dreiecke. Lassen wir nun alle Axiome der gewöhnlichen Euklidischen Geometrie bestehen mit Ausnahme des Axioms von der Dreieckskongruenz, indem wir vielmehr dieses durch das andere, weniger aussagende Axiom, daß in jedem Dreieck die Summe zweier Seiten größer als die dritte sein soll, ersetzen, so gelangen wir zu einer Geometrie, welche keine andere ist als diejenige, die Minkowski aufgestellt und zur Grundlage seiner geometrischen Untersuchungen gemacht hat. Diese *Minkowskische Geometrie* ist dann im wesentlichen durch folgende Festsetzungen charakterisiert:

1. Zwei Strecken heißen dann einander gleich, wenn man sie durch Parallelverschiebung des Raumes ineinander überführen kann.

2. Die Punkte, die von einem festen Punkte O gleichen Abstand haben, werden durch eine gewisse konvexe geschlossene Fläche des gewöhnlichen Euklidischen Raumes mit O als Mittelpunkt repräsentiert, so daß an Stelle der konzentrischen Kugeln der gewöhnlichen Euklidischen Geometrie ein System ineinander geschachtelter, durch Ähnlichkeitstransformation erzeugter konvexer Flächen tritt.

Insofern in der Minkowskischen Geometrie das Parallelenaxiom gilt, dagegen an Stelle des Axioms von der Dreieckskongruenz der gewöhn-

lichen Euklidischen Geometrie jenes weniger aussagende Axiom tritt, daß im Dreieck die Summe zweier Seiten die dritte übertrifft, ist die Minkowskische Geometrie eine der gewöhnlichen Euklidischen Geometrie nächststehende Geometrie, ebenso wie die Bolyai-Lobatschefskysche Geometrie, zu der sie ein Gegenstück bildet. Wie die Bolyai-Lobatschefskysche Geometrie in verschiedenen mathematischen Disziplinen, besonders in der Theorie der analytischen Funktionen mit linearen Transformationen in sich, die fruchtbarste Anwendung findet, so zeigt sich die Minkowskische Geometrie besonders für die Zahlentheorie von hervorragender Bedeutung.

Übertragen wir die eben angestellten geometrischen Überlegungen ins Analytische. In gewöhnlichen rechtwinkligen Koordinaten $x_1, \dots, x_n$ des n -dimensionalen Raumes kann die Oberfläche eines konvexen Körpers in der Gestalt

$$f(x_1, \dots, x_n) = 1$$

dargestellt werden, so daß f eine positive homogene (nicht notwendig rationale) Funktion ersten Grades bedeutet, deren wesentlichste Eigenschaft die ist, die durch die Funktionalungleichung

$$f(x_1 + y_1, \dots, x_n + y_n) \leq f(x_1, \dots, x_n) + f(y_1, \dots, y_n)$$

zum Ausdruck gebracht wird. Die Minkowskische Entfernung zwischen zwei Punkten $x_1, \dots, x_n$ und $y_1, \dots, y_n$ wird dann allgemein durch den Ausdruck

$$f(x_1 - y_1, \dots, x_n - y_n)$$

definiert. Die ursprünglich zugrunde gelegte Fläche f

$$f(x_1, \dots, x_n) = 1$$

heißt Eichfläche; sie ist das Minkowskische Analogon der Kugel im gewöhnlichen Euklidischen Raume.

Das Ausgangsbeispiel des Ellipsoides erhält man, wenn man hier für f die Funktion $\sqrt{F}$ nimmt, wo F die oben (S. X) erwähnte quadratische Form bedeutet.

Nun werde als Eichkörper ein konvexer Körper mit Mittelpunkt, d. h. ein solcher konvexer Körper genommen, der einen Punkt im Innern aufweist, in welchem alle hindurchgehenden Sehnen des Körpers halbiert werden. Dann gilt für die so definierte Minkowskische Entfernung der Satz, daß für die kleinste Entfernung zwischen zwei Gitterpunkten, d. h. für M , eine obere Schranke existiert, die allein vom Volumen des Eichkörpers abhängt; und zwar schließt man leicht, daß ein konvexer Körper mit einem Mittelpunkte in einem Punkte des Zahlengitters und vom Volumen 2^n immer noch mindestens zwei weitere Punkte des Zahlengitters, sei es im Innern, sei es auf der Begrenzung, enthalten muß.

Dieser Satz ist einer der anwendungsreichsten der Arithmetik; aus ihm leitet Minkowski seinen bekannten Determinantensatz ab, demzufolge

man in irgend n ganzen homogenen linearen Formen von n Variablen mit beliebigen reellen Koeffizienten und der Determinante 1 immer den Variablen solche ganzzahligen Werte, die nicht sämtlich Null sind, erteilen kann, daß dabei alle Formen absolute Beträge ≤ 1 erlangen; ferner die das Wesen der algebraischen Zahl tief berührende Tatsache, daß die Diskriminante eines algebraischen Zahlkörpers stets von ± 1 verschieden ist, d. h. daß es für einen algebraischen Zahlkörper stets wenigstens eine durch das Quadrat eines Primideals teilbare Primzahl, eine sogenannte Verzweigungszahl, gibt, analog wie in der Theorie der algebraischen Funktionen bekanntlich gezeigt wird, daß eine algebraische Funktion stets Verzweigungspunkte besitzen muß.

Aber der obige Satz vom Volumen des Eichkörpers, den ich einen der anwendungsreichsten der Arithmetik nannte, bildet doch nur das Anfangsglied einer Reihe weiterer auf geometrischer Anschauung fußender Schlußweisen von weittragender Bedeutung. So gelangt Minkowski durch eine sehr sinnreiche geometrische Überlegung, bei der der zugrunde gelegte konvexe Körper sukzessive nach bestimmten Vorschriften dilatiert wird, zu einer Erweiterung des ursprünglichen Satzes, die so lautet: Ist das Volumen des Eichkörpers gleich 2^n , so ist nicht nur, wie oben behauptet, die kleinste Minkowskische Entfernung ≤ 1 , sondern sogar das Produkt der n kleinsten Entfernungen, in n unabhängigen Richtungen genommen, fällt stets ≤ 1 aus. Die Endlichkeit der Klassenanzahl der positiven quadratischen Formen von n Variablen mit gegebener Determinante ist unter andern eine leichte Folge dieses allgemeinen Satzes.

Wie oben ausgeführt wurde, hat Minkowski für das Minimum einer quadratischen Form F von n Variablen mit der Determinante 1 mittels seiner geometrischen Methode eine obere, nur von n abhängige Schranke aufgestellt. Das genaue Minimum, d. h. der kleinste von Null verschiedene Wert, den F für ganzzahlige Variablen erlangt, ist notwendig noch eine Funktion der Koeffizienten der Form F ; lassen wir diese beliebig variieren, so jedoch, daß die Determinante beständig 1 bleibt, so können wir nach dem Maximum k_n der Minima aller dieser Formen fragen; dasselbe wird eine nur von n abhängige Zahl sein, welche jene obere Schranke ebenfalls nicht übersteigen kann. Durch völlig andere Hilfsmittel, aber ebenfalls ausgehend von einer geometrischen Betrachtung, bei der nunmehr der Begriff des Strahlenkörpers an Stelle des konvexen Körpers die wesentlichste Rolle spielt — Strahlenkörper ist ein Körper mit einem gewissen Punkte im Innern, der alle Strecken zwischen diesem Punkte und einem beliebigen Punkte des Körpers ganz enthält, so daß ein Strahlenkörper von einem gewissen Punkte aus diejenige Eigenschaft aufweist, welche bei einem konvexen Körper für jeden seiner Punkte erfüllt ist — gelangt

Minkowski für jenes Maximum k_n des Minimums der quadratischen Form F auch zu einer unteren Schranke. Ein überraschendes und für die Genauigkeit der Minkowskischen Methode zeugendes Resultat ist es, daß diese untere Schranke und die früher gefundene obere Schranke asymptotisch für $n = \infty$ ineinander fließen, so daß Minkowski die Limesgleichung

$$\lim_{n \rightarrow \infty} \frac{\log k_n}{\log n} = 1$$

aussprechen konnte.

Ch. Hermite, damals der Senior der französischen Mathematiker, hatte von Anbeginn die zahlentheoretischen Arbeiten Minkowskis mit höchstem Interesse und lebhaftester Freude verfolgt. Es ist rührend, wie rückhaltlos er die Vorzüge der Minkowskischen Methode gegenüber seinen eigenen Entwicklungen anerkennt, als Minkowski ihm die eben besprochenen Resultate mitteilt. „Au premier coup d'œil j'ai reconnu“, so schreibt Ch. Hermite in einem der an Minkowski gerichteten Briefe, „que vous avez été bien au delà de mes recherches en nous ouvrant dans le domaine arithmétique des voies toutes nouvelles.“ Und in einem zwei Jahre späteren Briefe vom November 1892 heißt es: „Je me sens rempli d'étonnement et de plaisir devant vos principes et vos résultats, ils m'ouvrent comme un monde arithmétique entièrement nouveau, où les questions fondamentales de notre science sont traitées avec un éclatant succès auquel tous les géomètres rendront hommage. Vous voulez bien, Monsieur, — et je vous en suis sincèrement reconnaissant — rapporter à mes anciennes recherches le point de départ de vos beaux travaux, mais vous les avez tant dépassées qu'elles ne gardent plus d'autre mérite que d'avoir ouvert la voie dans laquelle vous êtes entré.“

Hiernach nimmt es nicht Wunder, daß Hermites Begeisterung für die zahlentheoretischen Methoden Minkowskis keine Grenzen kannte, als die erste Lieferung seiner Geometrie der Zahlen 1896 erschien. „Je crois voir la terre promise“, so schreibt Hermite an Laugel, von dem er sich eine Übersetzung des Minkowskischen Buches zu seinem persönlichen Gebrauch anfertigen ließ. Und in der Tat, welche Fülle der verschiedenartigsten und tieflegendsten arithmetischen Wahrheiten werden in diesem Hauptwerke Minkowskis durch das geometrische Band gehalten und verknüpft! Die Theorie der Einheiten in den algebraischen Zahlkörpern, Sätze über die Ordnung einer endlichen Gruppe von homogenen linearen ganzzahligen Substitutionen und über die Zahl der Transformationen einer positiven quadratischen Form in sich, der Beweis für die Endlichkeit der Klassenanzahl von positiven quadratischen Formen mit gegebener Determinante, die Annäherung an beliebig viele reelle Größen durch rationale Zahlen mit den gleichen Nennern, die Theorie der Linearformen mit

ganzen komplexen Koeffizienten, Sätze über Minima von Potenzsummen linearer Formen, die Theorie der Kettenbrüche usw. bilden, von den schon vorhin aufgeführten Gegenständen abgesehen, die Themata des Minkowskischen Buches über die Geometrie der Zahlen.

Minkowski legte besonderen Wert auf die Darstellung, die er in seinem Buche der Theorie der gewöhnlichen Kettenbrüche hat zuteil werden lassen; er war der Meinung, daß durch seine geometrische Veranschaulichung erst das wahre Wesen des Kettenbruches enthüllt werde. In einer späteren Arbeit*) gelangt er, ebenfalls geleitet durch ein geometrisches Verfahren, welches in der sukzessiven Konstruktion von Parallelogrammen besteht, zu einer neuen Art von Kettenbruchentwicklung für eine beliebige reelle Zahl α . Diese Minkowskische Kettenbruchentwicklung ist so beschaffen, daß die dabei auftretenden Näherungsbrüche $\frac{x}{y}$ auch ohne Vermittlung des Kettenbruches direkt durch die Ungleichung

$$\left| \alpha - \frac{x}{y} \right| < \frac{1}{2} \frac{1}{y^2}$$

charakterisiert werden können; sie stellt demnach das bis dahin vermißte Analogon in der Größentheorie dar zu der in der Funktionentheorie üblichen Kettenbruchentwicklung, bei der ja ebenfalls die sämtlichen Näherungsbrüche, die der Kettenbruch einer Potenzreihe liefert, auch ohne den Kettenbruch unmittelbar definierbar sind.

Die Schlußlieferung von Minkowskis Geometrie der Zahlen ist nicht mehr erschienen, doch hat Minkowski den Stoff, den er für diese Lieferung plante, im wesentlichen in seinen späteren Abhandlungen zur Darstellung gebracht**).

Wenn wir uns diesen zuwenden, so haben wir vor allem eines Problems zu gedenken, dem Minkowski schon früh sein lebhaftes Interesse schenkte und auf welches er dann die in der ersten Lieferung seines Buches entwickelten Methoden mit sehr bemerkenswertem Erfolge anwandte***). Nach Lagrange fällt bekanntlich die Entwicklung einer reellen Zahl in einen Kettenbruch immer dann und nur dann periodisch aus, wenn die Zahl Wurzel einer quadratischen Gleichung mit rationalen Koeffizienten ist. Insofern dieser Satz ein notwendiges und hinreichendes

*) Ueber die Annäherung an eine reelle Größe durch rationale Zahlen. Mathematische Annalen, Bd. 54 (1901), S. 91—124. Diese Ges. Abhandlungen, B. I, S. 320—352.

**) Die posthum veröffentlichte „zweite Lieferung der Geometrie der Zahlen“ (Leipzig 1910) entspricht nicht der ursprünglich von Minkowski geplanten Schlußlieferung, sondern bringt vielmehr nur das fünfte Kapitel der ersten Lieferung zum Abschluß.

***) Ein Kriterium für die algebraischen Zahlen. Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen, mathematisch-physikalische Klasse, 1899, S. 64—88. Diese Ges. Abhandlungen, Bd. I, S. 293—315.

Kriterium für die quadratische Irrationalität enthält, lag es nahe, einen entsprechenden Satz für die algebraische Irrationalität beliebigen Grades n aufzustellen; doch waren alle bis dahin in dieser Richtung liegenden Versuche — ich erinnere an den Jacobischen Kettenbruchalgorithmus zur Entwicklung der kubischen Irrationalität, dessen Periodizität noch bis heute nicht festgestellt ist — vergeblich geblieben. Es gelang Minkowski zum ersten Male auf Grund sehr tiefliegender arithmetischer Sätze, zu deren Beweis seine geometrischen Methoden herangezogen werden, das gewünschte Kriterium für die algebraischen Zahlen beliebigen Grades n zu gewinnen. Der Minkowskische Algorithmus ist nicht ganz einfach; er besteht zunächst in einer Vorschrift, wie man aus der beliebig vorgelegten Zahl α in eindeutig bestimmter Weise eine Kette von gewissen linearen Substitutionen von n Variablen bestimmt und alsdann aus diesen gewisse lineare Formen ableitet: die Zahl α ist dann algebraisch vom Grade n , wenn die Kette niemals abbricht und zugleich alle jene unendlich vielen Formen aus einer endlichen Anzahl unter ihnen durch Multiplikation mit Faktoren entstehen.

In einer weiteren Untersuchung über die periodische Approximation algebraischer Zahlen*) beantwortet dann Minkowski insbesondere die Frage nach denjenigen algebraischen Zahlen α , für welche jene Substitutionen periodischen Charakter aufweisen, denen also in diesem Sinne genau die von Lagrange für die quadratische Irrationalität entdeckte Eigenschaft zukommt. Minkowski fand, daß die verlangte Periodizität außer für die quadratische Irrationalität nur noch in fünf ganz bestimmten Fällen stattfindet: nämlich im Falle $n = 3$, α komplex; ferner $n = 3$, α reell, während die zu α konjugierten Zahlen komplex sind; im Falle $n = 4$, wenn α nebst allen konjugierten Zahlen komplex ist, und endlich in je einem speziellen Fall bei $n = 4$ und $n = 6$.

Hatte Minkowski das ganze von ihm erschlossene Gebiet Geometrie der Zahlen genannt, weil er zu den Methoden, aus denen seine arithmetischen Sätze fließen, durch räumliche Anschauung geführt worden war, so blieb er auch bei der weiteren Erforschung dieses Gebietes stets dem Bestreben treu, durch engen Anschluß an die geometrischen Vorstellungen und Bilder die Fruchtbarkeit seiner Methoden zu zeigen; er wird nicht müde, durch originelle Modifikationen seine ursprünglichen Überlegungen zu vertiefen, die gefundenen arithmetischen Sätze zu vervollkommen und neue zu ersinnen.

So gelangt Minkowski**) zu einer gitterförmigen Bedeckung der Ebene

*) Ueber periodische Approximationen algebraischer Zahlen. Acta Mathematica, Bd. 26 (1902), S. 333—351. Diese Ges. Abhandlungen, Bd. I, S. 357—371.

**) Ueber die Annäherung an eine reelle Größe durch rationale Zahlen. Mathematische Annalen, Bd. 54 (1901), S. 108 ff. Diese Ges. Abhandlungen, Bd. I, S. 336 ff.

mit Parallelogrammen, bei der die ganze Ebene vollständig und andererseits keine Partie der Ebene mehr als zweifach überdeckt wird; diese Tatsache führt ihn unmittelbar zu einem Satze von Tschebyscheff über nichthomogene lineare diophantische Ungleichungen und zwar in einer allgemeineren und vollkommeneren Form, als derselbe von Tschebyscheff aufgestellt worden war.

Ferner wirft Minkowski die Frage auf*), unendlich viele untereinander kongruente und parallel orientierte Körper derart anzuordnen, daß sie, ohne einander zu durchdringen, sich so dicht als überhaupt möglich zusammenschließen, während ihre Schwerpunkte ein parallelepipedisches Punktsystem bilden. Wählt man für die Körper Kugeln, so zeigt sich dann, daß im Raume von drei Dimensionen zwar die bekannte tetraëdrale Anordnung von Kugeln die dichteste ist, daß aber in Räumen von höheren Dimensionen die dieser entsprechende tetraëdrale Anordnung keineswegs die dichteste Kugellagerung liefert. Das Problem der dichtesten Lagerung von Kugeln im n -dimensionalen Raum läuft auf die Bestimmung des Maximums k_n hinaus und hängt zusammen mit der Frage nach der Reduktion der positiven quadratischen Formen; diesem Problem wendet sich Minkowski in seiner zahlentheoretischen Abhandlung über den Diskontinuitätsbereich für arithmetische Äquivalenz**) noch einmal zu, es in vollendeter Form lösend, gleichsam als offensichtliches Wahrzeichen für die Leichtigkeit und Überlegenheit seiner gegenwärtigen mehr geometrischen Methoden im Vergleich zu dem Standpunkt seiner Jugendarbeiten.

Die Beweise der allgemeinen Sätze: der reduzierte Raum für die positiven quadratischen Formen von n Variablen ist eine konvexe Pyramide mit der Spitze im Nullpunkt, die von einer endlichen Anzahl durch diesen Punkt laufender Ebenen begrenzt wird; und: im Gebiet der positiv-definiten Formen grenzt der reduzierte Raum nur an eine endliche Anzahl von äquivalenten Räumen an; ferner die Berechnung des Volumens des reduzierten Raumes für alle Formen, deren Determinante eine gegebene Grenze nicht übersteigt, sowie die Anwendung hiervon auf die Bestimmung des asymptotischen Wertes der Klassenanzahl positiver quadratischer Formen sind die Glanzpunkte dieser letzten und inhaltreichsten zahlentheoretischen Abhandlung Minkowskis.

Von der Bedeutung der Zahlentheorie, wie sie in den Werken ihrer Heroen Fermat, Euler, Lagrange, Legendre, Gauß, Hermite, Dirichlet, Kummer, Jacobi und in deren begeisterten Aussprüchen sich wider-

*) Dichteste gitterförmige Lagerung kongruenter Körper. Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen, mathematisch - physikalische Klasse, 1904, S. 311—355. Diese Ges. Abhandlungen, Bd. II, S. 3—42.

**) Diskontinuitätsbereich für arithmetische Äquivalenz. Crelles Journal, Bd. 129 (1905), S. 220—274. Diese Ges. Abhandlungen, Bd. II, S. 53—100.

spiegelt, war Minkowski aufs tiefste durchdrungen; ihre Reize empfand er jederzeit aufs lebhafteste: war doch, was man an der Zahlentheorie rühmt, die Einfachheit ihrer Grundlagen, die Genauigkeit ihrer Begriffe und die Reinheit ihrer Wahrheiten ganz und gar zu seinem Wesen passend und seiner innersten Neigung am meisten zusagend. Wenn es zutrifft, daß nur ein enger Kreis von Mathematikern der Pflege der Zahlentheorie sich hingibt und so viele „von den eigenartigen, durch die Zahlentheorie ausgelösten Stimmungen kaum einen Hauch verspüren“: den Grund hierfür erblickt er darin, daß die Schöpfungen eines Gauß und der andern Großen zu erhaben sind. Und um in dieser gewaltigen Musik, wie er die Zahlentheorie nennt, für diejenigen, die nicht nur erbaut, sondern auch ergötzt sein wollen, die einschmeichelnden Melodien herauszuheben und so zu ihrem Genusse mehr anzulocken, dazu veröffentlichte er die Vorlesung, die er Winter 1903/04 in Göttingen gehalten hat, und in welcher er in leicht faßlicher Weise ohne die Voraussetzung besonderer Vorkenntnisse die wichtigsten Grundsätze der Geometrie der Zahlen und die einfachsten Anwendungen auf die Theorie der quadratischen Formen, auf die Zahlkörper und vor allem auf die Annäherung reeller und komplexer Größen durch rationale Zahlen auseinandersetzt. Das so entstandene Buch „*Diophantische Approximationen*“*) kann vorzüglich zur Einführung in die von Minkowski geschaffenen Methoden dienen.

Minkowski ist es zu danken, daß nach Hermites Tode die Führerrolle in der Zahlentheorie wieder in deutsche Hände zurückfiel und, wenn man überhaupt bei einer solchen Wissenschaft, wie es die Arithmetik ist, die Beteiligung der Nationen an den Fortschritten und Errungenschaften abwägen will: wesentlich durch Minkowskis Wirken ist es gekommen, daß heute im Reiche der Zahlen die bedingungslose und unbestrittene deutsche Vorherrschaft statthat.

Die Überzeugung von der tiefen Bedeutung des Begriffes eines konvexen Körpers, dessen Verwendung in der Zahlentheorie so erfolgreich gewesen war, hatte sich bei Minkowski immer mehr befestigt, und dieser Begriff bildet denn auch das Bindeglied zwischen denjenigen Arbeiten Minkowskis, die wesentlich zahlentheoretische Ziele im Auge haben, und seinen rein geometrischen Untersuchungen.

Das ursprüngliche Ziel, das Minkowski bei seinen rein geometrischen Untersuchungen im Auge hatte, war, die Begriffe Länge und Oberfläche mittels des Begriffes Volumen, „dieses elementarsten Begriffes der Analysis des Unendlichen“, zu erfassen**). In der Tat gelingt ihm diese Reduktion

*) *Diophantische Approximationen*. Eine Einführung in die Zahlentheorie. Leipzig 1907.

**) Volumen und Oberfläche. *Mathematische Annalen*, Bd. 57 (1903), S. 447—495. Diese Ges. Abhandlungen, Bd. II, S. 230—276.

durch ein einfaches Grenzverfahren. Ist etwa eine Kurve im Raume gegeben, so denkt sich Minkowski um jeden ihrer Punkte eine Kugel mit dem Radius r abgegrenzt. Das Volumen des so insgesamt in der Umgebung der Kurve abgegrenzten Bereiches nach Division durch den Inhalt des Kreises vom Radius r strebt in der Grenze für verschwindende Werte von r im allgemeinen einer Größe zu, die nunmehr als die Länge der Kurve eingeführt wird. Ähnlich kann der Begriff des Inhaltes einer Fläche eingeführt werden, und insbesondere die so entstehende Definition der Oberfläche ist es, durch die Minkowski zu einer wichtigen Verallgemeinerung des Begriffes der Oberfläche gelangt, indem er nämlich an Stelle von Kugeln beliebige einander ähnliche und ähnlich gelegene konvexe Körper verwendet — genau im Sinne der vorhin bei Besprechung der zahlentheoretischen Abhandlungen geschilderten Minkowskischen Geometrie.

Durch den Ausbau des Gedankens, die Kugel durch einen beliebigen Eichkörper zu ersetzen, gelangt Minkowski zu demjenigen Begriffe, der das Fundament seiner ganzen Theorie bildet, zu dem *Begriffe des gemischten Volumens* von irgend drei konvexen Körpern. Das gemischte Volumen von drei konvexen Körpern K_1, K_2, K_3 ist eine ganz bestimmte eindeutig aus denselben durch ein dreifaches Integral darzustellende Zahl V_{123} , die in das gewöhnliche Volumen eines Körpers übergeht, wenn man jene drei Körper miteinander identifiziert, die in die gewöhnliche Oberfläche eines Körpers übergeht, wenn man zwei von jenen drei Körpern miteinander identifiziert und den dritten gleich der Kugel mit dem Radius 1 nimmt und die endlich mit der totalen mittleren Krümmung der Oberfläche eines Körpers übereinstimmt, wenn man für zwei von jenen drei Körpern die Kugel mit dem Radius 1 wählt. So erscheint der Begriff des gemischten Volumens als der einfachste übergeordnete Begriff, der die Begriffe Volumen, Oberfläche, totale mittlere Krümmung als Spezialfälle enthält, und diese letzteren Begriffe sind damit in viel engeren Zusammenhang miteinander gebracht; steht doch deshalb auch von vornherein zu erwarten, daß wir auf diesem Standpunkte über das Verhältnis zwischen jenen Begriffen einen weit tieferen und allgemeineren Aufschluß erhalten, als bisher möglich war. Das Hauptergebnis, welches in dieser Hinsicht die Minkowskische Theorie liefert, gipfelt in der Ungleichung

$$V_{123}^2 \geq V_{122} V_{133},$$

einer Ungleichung, die lediglich quadratischen Charakter trägt, während beispielsweise der bekannte Satz, daß die Kugel unter allen Körpern gleicher Oberfläche das größte Volumen besitzt, für Volumen V und Oberfläche O eines beliebigen Körpers durch die kubische Ungleichung

$$36\pi V^2 \geq O^3$$

ausgedrückt wird. Diese kubische Ungleichung aber und somit insbesondere jener Satz über das Maximum des Kugelvolumens erscheint bei Minkowski als spezieller Ausfluß der genannten inhaltreicheren und einfacheren quadratischen Ungleichung; zugleich treten neben jenen Satz vom Maximum des Kugelvolumens eine ganze Reihe gleich wichtiger Sätze über die Kugel. Über das gemischte Volumen stellt Minkowski den allgemeinen Satz auf, daß, wenn man aus drei Körpern vom Volumen 1 das gemischte Volumen bildet, dieses stets ≥ 1 ist und nur dann gleich 1 wird, wenn die drei Körper miteinander identisch sind oder durch Translation miteinander zur Deckung gebracht werden können — ein Satz, der ebenfalls die in Rede stehende Maximaleigenschaft der Kugel als spezielle Folge mit enthält.

Zur analytischen Durchführung dieser Gedanken bedient sich Minkowski im wesentlichen der Methode der Ebenenkoordinaten. Die letzteren erscheinen in der Tat als das naturgemäße Hilfsmittel zur Darstellung der Minkowskischen Theorie; ist doch das Mischvolumen nichts Anderes als eine zweimalige Bildung der ersten Variation des gewöhnlichen Volumens, falls man dieses durch Ebenenkoordinaten ausdrückt.

Des weiteren beschäftigt sich Minkowski mit dem einfachen und elementaren Begriffe des konvexen Polyeders und weiß diesem vielbehandelten Gegenstande neue und fruchtbare Seiten abzugewinnen. Sein grundlegender Satz sagt aus, daß ein konvexes Polyeder stets durch die Richtungen der Normalen und die Inhalte seiner Seitenflächen bis auf eine Translation eindeutig bestimmt wird. Aus diesem Satze leitet Minkowski durch Grenzübergang das merkwürdige Theorem ab, wonach es immer eine und nur eine geschlossene konvexe Fläche gibt, für die die Gaußsche Krümmung als stetige Funktion der Richtungskosinusse ihrer Normalen vorgeschrieben ist. Indem hierbei Minkowski die Krümmung — unmittelbar an die ursprüngliche Betrachtungsweise von Gauß anschließend — durch eine Integralforderung definiert, vermeidet er es, die Existenz der zweiten Ableitungen der die Fläche definierenden Funktion vorauszusetzen, und erreicht eben dadurch jene größtmögliche Einfachheit und Allgemeinheit in der Fassung und Entwicklung des Theorems.

Das Minkowskische Problem der Bestimmung der geschlossenen konvexen Flächen mit vorgeschriebener Gaußscher Krümmung ist wesentlich identisch mit dem Problem der Integration einer gewissen *partiellen Differentialgleichung vom Monge-Ampèreschen Typus*; so kommt es, daß die ursprünglich rein geometrische, auf dem Begriff des konvexen Körpers beruhende Methode Minkowskis zugleich für die Theorie der Integration gewisser nichtlinearer partieller Differentialgleichungen bis dahin unbekannte Fragestellungen und aussichtsreiche Angriffspunkte liefert.

Endlich werde noch eines kleinen Vortrages*) von Minkowski Erwähnung getan, den er vor seiner Übersiedelung nach Göttingen in der hiesigen mathematischen Gesellschaft gehalten hat und der bisher nur in einer russischen Übersetzung publiziert worden ist; derselbe enthält einen Satz von elementarem Charakter, wonach die Körper, deren Breite konstant d. h. in jeder Richtung genommen die nämliche ist, und andererseits die Körper konstanten Umfanges miteinander identisch sind; dabei ist unter Umfang der Umfang des Querschnittes des in irgendeiner Richtung dem Körper umschriebenen Zylinders zu verstehen.

Sein Interesse für die physikalische Wissenschaft hat Minkowski frühzeitig bekundet. Schon in den ersten Jahren seiner Privatdozentenzeit in Bonn beschäftigte er sich mit theoretischen Untersuchungen über *Hydrodynamik*. Helmholtz legte 1888 in der Akademie der Wissenschaften zu Berlin eine Arbeit**) von Minkowski über das Problem der kräftefreien Bewegung eines beliebigen starren Körpers in einer reibungslosen inkompressiblen Flüssigkeit vor. Um die Bewegung des Körpers völlig zu kennzeichnen, ist die Bestimmung von sechs unbekannten Funktionen der Zeit erforderlich. Das wichtigste Resultat von Minkowski besteht nun in der Reduktion des ursprünglich durch das Hamiltonsche Prinzip gelieferten Variationsproblems auf ein Variationsproblem, welches nur zwei unbekannte Funktionen der Zeit enthält.

Die Ferienzeiten während der Bonner Jahre verlebte Minkowski in der Regel in Königsberg, dem Wohnorte seiner Familie, wo er dann mit Hurwitz und mir fast täglich zusammenkam, meist auf Spaziergängen in der Königsberger Umgebung. Einmal, Weihnachten 1890, blieb Minkowski in Bonn; auf mein Zureden nach Königsberg zu kommen, stellte er sich in einem launigen Briefe als einen physikalisch völlig Durchseuchten hin, der erst eine zehntägige Quarantäne durchmachen müßte, ehe Hurwitz und ich ihn in Königsberg als mathematisch rein zu unseren Spaziergängen zulassen würden. „Ich habe mich“, so fährt Minkowski in seinem Briefe fort, „ganz der Magie, wollte sagen der Physik ergeben. Ich habe meine praktischen Übungen im physikalischen Institut, zu Hause studiere ich Thomson, Helmholtz und Konsorten; ja von Ende nächster Woche an arbeite ich sogar an einigen Tagen der Woche in blauem Kittel in einem Institut zur Herstellung physikalischer Instrumente, also ein Praktikus schändlichster Sorte.“ Von Heinrich Hertz in Bonn fühlte sich Minkowski

*) Ueber die Körper konstanter Breite. Moskau, Mathematische Sammlung (Matematičeskij Sbornik), Bd. 25 (1906), S. 505—508. Diese Ges. Abhandlungen, Bd. II, S. 277—279.

**) Ueber die Bewegung eines festen Körpers in einer Flüssigkeit. Sitzungsberichte der Berliner Akademie, 1888, S. 1095—1110. Diese Ges. Abhandlungen, Bd. II, S. 283—297.

stark angezogen; er äußerte, daß er, wenn Hertz am Leben geblieben wäre, sich schon damals mehr der Physik zugewandt hätte.

August 1892 war Minkowski zum außerordentlichen Professor in der philosophischen Fakultät zu Bonn ernannt worden. April 1894 ermöglichte auf Minkowskis und meinen dringenden Wunsch der damalige Ministerialrat Althoff, der Scharfblickende, in dem Minkowski sehr frühzeitig einen Gönner und Bewunderer gefunden hatte, die Versetzung Minkowskis nach Königsberg, und ein Jahr später wurde Minkowski dann in Königsberg mein Nachfolger im dortigen Ordinariat für Mathematik. Aus diesem Amte schied er Oktober 1896, um einem Rufe als Professor für Mathematik an das Eidgenössische Polytechnikum in Zürich zu folgen. Dort verheiratete er sich im Jahre 1897 mit Auguste Adler aus Straßburg i. E. In Zürich blieb er bis zum Herbst 1902. Da war es wiederum Althoff, der Minkowski auf den für seine Wirksamkeit angemessensten Boden verpflanzte; mit einer Kühnheit, wie sie vielleicht in der Geschichte der Verwaltung der Preußischen Universitäten beispiellos dasteht, schuf Althoff aus nichts hier in Göttingen eine neue ordentliche Professur, und dieser Tat Althoffs danken wir es, daß seit Herbst 1902 Minkowski der unsrige gewesen ist. Bereits Oktober 1901 hatte ihn unsere Gesellschaft zu ihrem korrespondierenden Mitgliede in der mathematisch-physikalischen Klasse gewählt.

Als Frucht der vielseitigen theoretisch-physikalischen Studien, die Minkowski auch in Zürich betrieben hatte und in Göttingen fortsetzte, ist der Enzyklopädieartikel über *Kapillarität**) anzusehen, in welchem er in wahrhaft musterhafter Weise in aller Kürze, dem beschränkten Raum entsprechend, die sämtlichen theoretischen Gesichtspunkte dieses Kapitels der Physik auseinandersetzt und die schwierigen mathematischen Grundlagen, insbesondere soweit sie die Variationsrechnung betreffen, in origineller, zum Teil ganz neuer Form entwickelt.

Aber am nachhaltigsten fesselten Minkowski die modernen elektrodynamischen Theorien, die er mehrere Semester hindurch mit mir gemeinsam betrieb, insbesondere in Vorträgen, zu denen das von ihm und mir geleitete Seminar Anlaß bot. Die letzten Schöpfungen Minkowskis entsprangen diesen Studien, denen er mit großem Eifer oblag; hatte er doch für die nächsten Semester Vorlesungen und Seminar über Elektronentheorie geplant.

H. A. Lorentz hat zuerst erkannt, daß die Grundgleichungen der Elektrodynamik für den reinen Äther die Eigenschaft der Invarianz gegenüber denjenigen gleichzeitigen Transformationen der Raumkoordinaten x, y, z und

*) Enzyklopädie der mathematischen Wissenschaften, Bd. V 1, Heft 4, S. 558—613. Diese Ges. Abhandlungen, Bd. II, S. 298—351.

des Zeitparameters t besitzen, die — falls man die Lichtgeschwindigkeit gleich 1 nimmt — den Ausdruck $x^2 + y^2 + z^2 - t^2$ in sich überführen. Im Zusammenhang mit dieser rein mathematischen Tatsache und in der Absicht, davon Rechenschaft zu geben, daß eine relative Bewegung der Erde gegen den Lichtäther nicht wahrgenommen wird, war jener scharfsinnige Forscher in kühnem Gedankenfluge zu der Einsicht gelangt, daß der Begriff des starren Körpers in dem bisherigen Sinne nicht aufrecht zu erhalten sei, sondern in der Weise modifiziert werden müsse, daß Elektrizität und Materie, sofern sie eine Bewegung von der Geschwindigkeit v besitzen, in Richtung dieser Bewegung eine Verkürzung ihrer Ausdehnung erfahren und zwar im Verhältnis $1:\sqrt{1-v^2}$. Daß eine weitere Konsequenz dieser Idee eine neuartige Auffassung des Zeitbegriffes ist, und insbesondere alle den Lorentz-Transformationen entsprechenden Bezugssysteme zur Einführung eines Zeitparameters gleichberechtigt sind, dies erkannt zu haben, ist das Verdienst des Physikers Einstein.

Die Ideenbildungen von Lorentz und Einstein, die man unter dem Namen des Relativitätsprinzipes zusammenfaßt, waren es, die Minkowski die Anregung zu seinen wichtigen und auch in weiteren Kreisen bekannt gewordenen *elektrodynamischen Untersuchungen* gaben. Minkowski*) legte sofort jener mathematischen Tatsache der Invarianz der elektrodynamischen Grundgleichungen gegenüber den Lorentz-Transformationen die allgemeinste und weitgehendste Bedeutung bei, indem er diese Invarianz als eine Eigenschaft auffaßte, die überhaupt allen Naturgesetzen zukomme, ja daß sie nichts Anderes als eine schon in den Begriffen Raum und Zeit selbst enthaltene und diese beiden Begriffe gegenseitig verkettende und miteinander verschmelzende Eigenschaft sei. Auch dem Nicht-Naturforscher ist die Tatsache geläufig, daß die Naturgesetze von der Orientierung im Raume sowie von der Zeit unabhängig sind, und ferner lehrt die gewöhnliche Mechanik, daß, wenn ein System sich bewegt, stets auch diejenige Bewegung statthaben kann, bei welcher die Geschwindigkeitsvektoren sämtlicher materieller Punkte je um einen konstanten Vektor vermehrt sind: darüber hinaus behauptet nun nach Minkowski das Relativitätsprinzip — oder, wie es Minkowski später nennt, das *Weltpostulat* —, daß die Naturgesetze in einem noch viel höheren Sinne von Raum und Zeit unabhängig, nämlich invariant gegenüber allen Lorentz-Transformationen sind. Indem nun durch die Lorentz-Transformationen gewisse Abänderungen des Zeitparameters zugelassen werden, die nicht bloß auf eine veränderte Wahl des Zeitanfanges hinauslaufen, fällt konsequenterweise überhaupt der Be-

*) Die Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern. Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen, mathematisch-physikalische Klasse, 1908, S. 53—111. Diese Ges. Abhandlungen, Bd. II, S. 352—404.

griff der Gleichzeitigkeit zweier Ereignisse als an sich existierend. Nur weil wir gewohnt sind, ein bestimmtes Bezugssystem für Raum und Zeit stark approximativ eindeutig zu wählen, halten wir den Begriff der Gleichzeitigkeit für einen absoluten — ungefähr wie Wesen, gebannt an eine enge Umgebung eines Punktes auf einer Kugeloberfläche, darauf verfallen könnten, die Kugel sei ein geometrisches Gebilde, an welchem ein Durchmesser an sich ausgezeichnet ist. Tatsächlich ist die Sachlage die, daß stets zwei Ereignisse, die an zwei Orten zu zwei verschiedenen Zeiten stattfinden, als gleichzeitig aufgefaßt werden können, sobald die Zeitdifferenz kleiner als die Entfernung beider Orte, d. h. diejenige Zeit ausfällt, die das Licht braucht, um von dem einen Orte zu dem andern zu gelangen. Ähnlich verhält es sich mit drei Ereignissen zu drei verschiedenen Zeiten, die ebenfalls als gleichzeitig stattfindend aufgefaßt werden können, sobald gewisse Ungleichheiten zwischen den Raum- und Zeitparametern erfüllt sind. Erst durch vier Ereignisse ist im allgemeinen das Bezugssystem von Raum und Zeit eindeutig festgelegt. — „Von Stund an sollen Raum für sich und Zeit für sich völlig zu Schatten herabsinken, und nur noch eine Art Union der beiden soll Selbständigkeit bewahren.“ So bekannte sich Minkowski eingangs des eindrucksvollen Vortrages*), den er auf der vorjährigen Naturforscherversammlung zu Köln vor einer zahlreichen, ihm mit größter Aufmerksamkeit folgenden Zuhörerschaft, bestehend aus Mathematikern, Physikern und Philosophen, gehalten hat.

Um die in Rede stehende Invarianz der Naturgesetze richtig zu verstehen, ersetze man sowohl die Raum- und Zeitparameter x, y, z, t , wie auch diejenigen Größen, die in den die Naturgesetze ausdrückenden Gleichungen als Funktionen von x, y, z, t auftreten, durch die entsprechend linear transformierten Größen: dann müssen die erhaltenen Gleichungen die nämliche Form für die neuen Größen in den neuen Veränderlichen aufweisen. Beispielsweise sind im Falle der elektrodynamischen Grundgleichungen die mit der Dichte multiplizierten Geschwindigkeitskomponenten u, v, w zusammen mit der Dichte ρ als vier Größen anzusehen, die in gleicher Weise mit den Variablen x, y, z, t transformiert werden; die Vektorenpaare dagegen, der elektrische und der magnetische Vektor einerseits und die elektrische und magnetische Erregung andererseits, sind als je sechs Größen anzusehen, die wie die sechs zweireihigen Determinanten einer Matrix zweier Raumzeitpunkte, d. h. etwa wie die Plückerschen Linienkoordinaten sich transformieren. Da demnach bei diesen Transformationen eine Vermischung von Geschwindigkeiten und Dichte und

*) Raum und Zeit. Physikalische Zeitschrift, 10. Jahrgang, Nr. 3 (1909), S. 104—111; Jahresberichte der Deutschen Mathematiker-Vereinigung, Bd. 18 (1909), S. 75—88. Diese Ges. Abhandlungen, Bd. II, S. 431—444.

ebenso von elektrischen und magnetischen Vektoren stattfindet, so ist absolut genommen eine Festlegung von Geschwindigkeit und Dichte der Substanz, sowie der elektrischen und magnetischen Vektoren nicht möglich; diese Begriffe hängen vielmehr ebenfalls wesentlich von der Wahl des Bezugssystems für x, y, z, t ab.

Minkowski wendet nun das eben gekennzeichnete und von ihm mathematisch präzierte Weltpostulat — und darin erblicke ich seine bedeutendste positive Leistung auf diesem Gebiete — dazu an, um die elektrodynamischen Grundgleichungen für bewegte Materie, deren definitive Form unter den Physikern außerordentlich strittig war, herzuleiten. Dazu sind nur drei sehr einfache Grundannahmen nötig: nämlich

1) die Annahme, daß die Geschwindigkeit der Materie stets und an allen Orten kleiner als 1 d. h. als die Lichtgeschwindigkeit ist;

2) das Axiom, daß, wenn an einer einzelnen Stelle die Materie in einem Momente ruht — die Umgebung mag in irgendwelcher Bewegung begriffen sein — dann für jenen „Raumzeitpunkt“ zwischen den magnetischen und elektrischen Vektoren und deren Ableitungen nach x, y, z, t genau die nämlichen Beziehungen statthaben, die zu gelten hätten, falls alle Materie ruhte;

3) die Annahme der von niemand bestrittenen elektrodynamischen Grundgleichungen für ruhende Materie.

Die elektrodynamischen Grundgleichungen, die Minkowski auf diesem Wege erhält*), lassen, was Durchsichtigkeit und Einheitlichkeit betrifft, nichts zu wünschen übrig; sie stimmen mit den bisherigen Beobachtungen überein, weichen indes in mannigfaltiger Weise von den bis dahin gebrauchten, von Lorentz und Cohn aufgestellten Gleichungen ab, indem diese keineswegs das Weltpostulat genau erfüllen. Die *Minkowskischen elektrodynamischen Grundgleichungen* sind eine notwendige Folgerung des Weltpostulates — sie sind von derselben Gewißheit wie dieses.

Immer mehr und mehr befestigte sich Minkowski in der Überzeugung von der allgemeinen Gültigkeit und der eminenten Fruchtbarkeit und Tragweite seines Weltpostulats und — die wunderbaren, vielverheißenden Ideen von M. Planck über die Dynamik bewegter Systeme bestärkten ihn darin — von der Notwendigkeit einer Reform der gesamten Physik nach Maßgabe dieses Postulats.

Was die Mechanik betrifft, so gelangte Minkowski durch Einführung

*) Mit der Ausarbeitung einer Ableitung dieser Gleichungen auf Grund der Vorstellungen der Elektronentheorie war Minkowski in den letzten Wochen seines Lebens beschäftigt. Unter Benutzung der nachgelassenen Papiere ist eine solche Herleitung in Minkowskis Sinne von Herrn M. Born (Mathematische Annalen, Bd. 68 (1910), S. 526—551; diese Ges. Abhandlungen, Bd. II, S. 405—430) durchgeführt worden.

des Begriffs der Eigenzeit eines materiellen Punktes zu einem gewissen System modifizierter Newtonscher Bewegungsgleichungen, bestehend aus vier Gleichungen, von denen die drei ersten in die gewöhnlichen Newtonschen Gleichungen übergehen, wenn man die Lichtgeschwindigkeit c unendlich werden läßt, während die vierte eine Folge der drei ersten ist und den Satz von der Erhaltung der Energie ausspricht. In dieser dem Weltpostulat gemäß reformierten Mechanik fallen die Disharmonien zwischen der Newtonschen Mechanik und der modernen Elektrodynamik von selbst weg. Aber die Minkowskische Untersuchung führt darüber hinaus zu der prinzipiell interessanten Tatsache, daß auf Grund des Weltpostulates die vollständigen Bewegungsgesetze allein aus dem Satz von der Erhaltung der Energie ableitbar sind.

Ferner zeigte Minkowski, wie das Newtonsche Gravitationsgesetz zu modifizieren sei, damit es dem Weltpostulat genügt. Das *Minkowskische Gravitationsgesetz* verknüpft mit der *Minkowskischen Mechanik* ist nicht weniger geeignet, die astronomischen Beobachtungen zu erklären als das Newtonsche Gravitationsgesetz verknüpft mit der Newtonschen Mechanik. Dabei bedeutet die Minkowskische Formulierung eine Fortpflanzung der Gravitation mit Lichtgeschwindigkeit — was unserer heutigen Anschauungsweise über Fernwirkung weit besser entspricht als die alte Newtonsche Momentanwirkung.

Als Beleg dafür, wie die Minkowskische Betrachtungsweise, die sich stets in der vierdimensionalen Raum-Zeitmannigfaltigkeit x, y, z, t — Welt genannt — bewegt, erst imstande ist, die innere Einfachheit und den wahren Kern der Naturgesetze zu enthüllen, sei nur noch auf den wunderbar durchsichtigen, von Minkowski angegebenen Ausdruck für die so äußerst komplizierte ponderomotorische Wirkung zweier bewegter elektrischer Teilchen hingewiesen.

Damit ist die Würdigung der hauptsächlichsten Ergebnisse der Publikationen Minkowskis beendet; aber die wissenschaftliche Wirksamkeit seiner Person ist durch die zur Veröffentlichung gelangten Schriften keineswegs erschöpft. Nach welchen Richtungen weiterhin und in welchem Sinne sich diese Wirksamkeit Minkowskis vornehmlich erstreckte, bedarf noch einer kurzen Darlegung, da erst dann die volle Bedeutung Minkowskis für die Entwicklung der Mathematik der Gegenwart sich erkennen läßt.

Zunächst gedenke ich der Stellungnahme Minkowskis gegenüber derjenigen mathematischen Disziplin, welche heute eine hervorragende Rolle in unserer Wissenschaft einnimmt und ihren gewaltigen Einfluß auf alle Gebiete der Mathematik ausströmt, nämlich der Mengentheorie. Diese von Georg Cantor zuerst in fruchtbarer Weise in Angriff genommene und

durch kühne Ideen zu gewaltiger Höhe geführte Lehre wurde damals von dem im Gebiet der Zahlentheorie maßgebenden Mathematiker Kronecker aufs entschiedenste bekämpft. Obwohl Minkowski in Berlin bei Kronecker studiert hatte und sich dem mächtigen Einfluß, den dieser in der Zahlentheorie ausübte, willig hingab: die Vorurteile, von denen Kronecker befangen war, durchschaute er frühzeitig; er war der erste Mathematiker unserer Generation — und ich habe ihn darin nach Kräften unterstützt —, der die hohe Bedeutung der Cantorschen Theorie erkannte und zur Geltung zu bringen suchte. „Die spätere Geschichte“, so führt Minkowski in einem in Königsberg gehaltenen Vortrag über das Aktual-Unendliche in der Natur aus, „wird Cantor als einen der tiefstinnigsten Mathematiker dieser Zeit bezeichnen; es ist sehr zu bedauern, daß eine nicht auf sachlichen Gründen allein beruhende Opposition, die von einem sehr angesehenen Mathematiker“ — gemeint ist eben Kronecker — „ausging, Cantor die Freude an seinen wissenschaftlichen Forschungen trüben konnte.“ Minkowski verehrte in Cantor den originellsten zeitgenössischen Mathematiker zu einer Zeit, als in damals maßgebenden mathematischen Kreisen der Name Cantor geradezu verpönt war und man in Cantors transfiniten Zahlen lediglich schädliche Hirngespinnste erblickte. Minkowski äußerte wohl, daß Cantors Name noch genannt werden würde, wenn man die heute — weil sie modisch sind — im Vordergrund stehenden Mathematiker längst vergessen hat. Der Umstand, daß ein Mann wie Minkowski, der das exakte Schließen in der Mathematik gewissermaßen verkörperte und dessen Sinn für echte Zahlentheorie über allem Zweifel war, so urteilte, ist der Verbreitung der Cantorschen Theorie, „dieser ursprünglichen Schöpfung genialer Intuition und spezifischen mathematischen Denkens“, wie sie mit Recht kürzlich ein jüngerer Mathematiker genannt hat, sehr zustatten gekommen.

Minkowski hat stets danach gestrebt, nicht nur über die Methoden der reinen Mathematik die Herrschaft zu erlangen, sondern auch den wesentlichen Inhalt aller derjenigen Wissensgebiete sich anzueignen, in denen die Mathematik als Hilfswissenschaft eine entscheidende Rolle zu spielen berufen ist. Wie tief er dann in solche Wissensgebiete, die seinem eigentlichen Arbeitsfelde fern lagen, eindrang und wie kritisch auch hier sein Blick war, zeigen die mannigfachen Vorträge, die er bei verschiedenen Anlässen, namentlich in unserer mathematischen Gesellschaft, gehalten hat, sowie seine Universitätsvorlesungen. Zumal in Göttingen hat Minkowski außer den üblichen Vorlesungen eine große Anzahl von Spezialvorlesungen über die verschiedensten Gegenstände gehalten, z. B. über Linien- und Kugelgeometrie, Analysis situs, automorphe Funktionen, Invariantentheorie, Wärmestrahlung und Wahrscheinlichkeitsrechnung. Diese

Vorlesungen waren stets klar durchdacht und fein geformt; ihr Ziel war, die Ergebnisse neuester Forschung kritisch zu sichten, auf die einfachste Form zu bringen und alsdann in Verbindung mit den alten Sätzen der Theorie einheitlich zur Darstellung zu bringen. Wie sehr es ihm dabei gelang, auch den schwerfälligeren Zuhörern die Wege zu ebnen und die reiferen ganz für sich zu gewinnen, beweist der steigende Zuspruch, dessen sich diese Vorlesungen in Göttingen erfreuten. Besonders verstand er es, in höheren Vorlesungen junge Mathematiker zu eigenen Forschungen anzuregen. Unter den Dissertationen, die seiner Anregung zu verdanken sind, seien nur die von L. Kollros, *Un algorithme pour l'approximation simultanée de deux grandeurs* (1905), und E. Swift, *Über die Form und Stabilität gewisser Flüssigkeitstropfen* (1907), genannt, deren wertvolle Resultate in weiteren Fachkreisen bekannt geworden sind.

Daß Minkowski auch Nichtfachleuten durch die Heranziehung treffender Gleichnisse und anschaulicher Bilder über schwierige mathematische Gegenstände vorzutragen und in ihnen eine Vorstellung von der Größe und Erhabenheit unserer Wissenschaft zu erwecken wußte, zeigt am besten die Rede, die er in der Festsitzung der Göttinger mathematischen Gesellschaft zur hundertjährigen Wiederkehr des Geburtstages von Dirichlet gehalten hat*). Die begeisterten und klaren Ausführungen, die dort Minkowski über den Charakter der Zahlentheorie, ihre Bedeutung und ihre Stellung zu anderen Disziplinen machte, beruhen auf einer tiefen Erfassung des Wesens der Zahlentheorie und sind das Beste, was je über diese wunderbarste Schöpfung menschlichen Geistes gesagt worden ist. Hierfür sei das Zeugnis desjenigen Mathematikers angerufen, der als Schüler von Dirichlet ein kompetentes Urteil hat, und den wir heute im In- und Auslande als den Senior der Mathematiker, als den einzigen lebenden Heros aus der größten Epoche der Zahlentheorie verehren dürfen. „Ich habe Ihren Vortrag“, so schrieb Richard Dedekind an Minkowski, „mit größtem Genuß fünfmal und noch viel öfter durchgelesen und bin besonders von der großen historischen Auffassung ergriffen, mit der Ihr Vortrag die tiefsten Gedanken unserer Wissenschaft deutlich erfaßt und in ihrer Entwicklung verfolgt“.

Trotz seiner milden Denkart war Minkowski im Grunde kritisch, er erkannte leicht die Schwächen einer Beweisführung oder einer Ideenbildung und legte im allgemeinen auch an die Arbeiten anderer einen strengen Maßstab an. Er unterschied scharf zwischen oberflächlichen und soliden Mathematikern. Von einer guten mathematischen Arbeit ver-

*) P. G. Lejeune Dirichlet und seine Bedeutung für die heutige Mathematik. Jahresbericht der Deutschen Mathematiker-Vereinigung, Bd. 14 (1905), S. 149—163. Diese Ges. Abhandlungen, Bd. II, S. 447—461.

langte er, daß in ihr eine klar gestellte und des Interesses werthe Frage gelöst werde.

So sehr er von echter Bescheidenheit war und mit seiner Person gern im Hintergrunde blieb, war er doch von der innersten Überzeugung getragen, daß vieles von dem, was er schuf, die Arbeiten anderer zeitgenössischer Autoren überleben und einst zur allgemeinen Anerkennung gelangen würde. Den von ihm gefundenen Satz von der Lösbarkeit linearer Ungleichungen mit der Determinante 1, seinen Beweis für die Existenz von Verzweigungszahlen im Zahlkörper oder die Reduktion der kubischen Ungleichung, die die vorhin genannte Maximaleigenschaft der Kugel ausdrückt, auf eine quadratische Ungleichung stellte er wohl innerlich selbst den besten Leistungen der mathematischen Klassiker auf dem Gebiet der Zahlentheorie und Geometrie gleichwertig an die Seite.

Man müsse fleißig sein, das Leben sei ja so kurz, äußerte er wohl. Und in der Tat, die Wissenschaft begleitete ihn überall, sie war ihm zu jeder Zeit interessant und ermüdete ihn an keinem Ort, sei es auf einem Ausflug, in der Sommerfrische oder in der Bildergalerie, in dem Eisenbahncoupé oder auf dem Großstadtpflaster.

Noch in den letzten Nächten, die er zu Hause zubrachte, beschäftigte ihn die Formung der Worte in seinem Kölner Vortrage, und er überlegte, welche Wendung dem naiven Sprachgefühl besser entspräche. Das war charakteristisch für ihn: er strebte zuerst nach Einfachheit und Klarheit des Gedankens — Dirichlet und Hermite waren darin seine Vorbilder —, dann bemühte er sich, dem Gedanken auch eine vollkommene Darstellung zu geben. Er war von großer Genauigkeit und einer ins kleinste Detail gehenden Eigenheit, was die Wahl der Bezeichnungen und der Buchstaben betraf, eine Genauigkeit, die — freilich wie bei Minkowski gepaart mit einem aufs Große gerichteten Blick — dem rechten Forscher stets eigen ist, und die wir heute bedauerlicherweise seltener werden sehen. Auch sonst, wenn er im kleineren Kreise über einen wissenschaftlichen Gegenstand sprach, legte er auf die Form und den Ausdruck Wert, und besonders in unserer mathematischen Gesellschaft verfehlte er selten, seinem Vortrage einige wohl überlegte, die Zuhörer anregende Bemerkungen vorzuschicken.

Frei von aller vorgefaßten Meinung und von aller Einseitigkeit zeigte er auch für die entferntesten Anwendungen der Mathematik Interesse — immer der Meinung, daß diese auch der reinen Wissenschaft schließlich zum Vorteil dienen würden. So nahm er auch an den Sitzungen der Göttinger Vereinigung für angewandte Mathematik und Physik theil.

Er besaß eine scharfe Beobachtungsgabe auch für Dinge, die nicht

seine Wissenschaft betrafen. Wie er denn überhaupt für alles, was Menschen bewegt — von der Politik bis zum Theater — Verständnis, nicht selten Eifer und Lebhaftigkeit bekundete. Dem Fernerstehenden schien es mitunter bei dem im allgemeinen ruhigen Temperament Minkowskis, als schenke er einer Sache wenig Interesse: oft fiel gerade dann von Minkowskis Seite eine Bemerkung, die den Kern der Sache traf, oder er hatte gar ein Zitat aus Faust bereit, den er vollständig auswendig konnte. Noch in der letzten arbeitsreichsten Zeit seines Lebens liebte er es, seinen Kindern Gedichte von Goethe und Schiller auswendig vorzutragen — mit der Begeisterung, die ihm aus seiner Jugendzeit frisch geblieben war.

Für seine Person war er äußerst einfach und anspruchslos, mehr bedacht auf das Wohlergehen seiner Angehörigen als auf sein eigenes.

Er war von unentwegtem Optimismus, stets überzeugt, daß das Gute und Richtige zum schließlichen Siege gelangen würde. Für junge heranwachsende Mathematiker hatte er viel persönliches Interesse und sah sie häufig bei sich im Hause; er sprach sich bisweilen überschwenglich über die Kenntnisse und den Fleiß einzelner unter ihnen aus und setzte große Hoffnungen auf ihre Zukunft.

Seit meiner ersten Studienzeit war mir Minkowski der beste und zuverlässigste Freund, der an mir hing mit der ganzen ihm eigenen Tiefe und Treue. Unsere Wissenschaft, die uns das liebste war, hatte uns zusammengeführt; sie erschien uns wie ein blühender Garten; in diesem Garten gibt es geebnete Wege, auf denen man mühelos genießt, indem man sich umschaute, zumal an der Seite eines Gleichempfindenden. Gern suchten wir aber auch verborgene Pfade auf und entdeckten manche neue, uns schön dünkende Aussicht, und wenn der eine dem andern sie zeigte und wir sie gemeinsam bewunderten, war unsere Freude vollkommen.

Sein stiller Sinn stand nicht nach äußeren Zeichen der Anerkennung; doch empfand er eine lebhaftige Genugtuung, wenn mir eine solche zuteil wurde. Allem, was mich betraf, brachte er sein stets gleichbleibendes Interesse und seine herzlichste Teilnahme entgegen. Zumal die kleine Stadt hier erleichterte unsern Verkehr: ein Telephonruf zur Vermittlung einer Verabredung oder ein paar Schritte über die Straße und ein Steinchen an die klirrende Scheibe des kleinen Eckfensters seiner Arbeitsstube — und er war da, zu jeder mathematischen oder nichtmathematischen Unternehmung bereit.

Noch auf der Krankenbahre liegend — todeswund — galten seine Gedanken dem Bedauern, daß er in der nächsten Stunde des Seminars, in der ich meine Lösung des Waringschen Problems vortragen wollte, nicht zugegen sein könne. Seinem Andenken darum habe ich meine die Lösung

enthaltende Abhandlung gewidmet, die erste, von deren Inhalt er keine Kenntnis mehr genommen hat und über deren Korrekturbogen sein sicheres Auge nicht gegelitten ist.

Er war mir ein Geschenk des Himmels, wie es nur selten jemand zuteil wird, und ich muß dankbar sein, daß ich es so lange besaß.

Jeder, der ihm näher stand, empfand die Harmonie seiner Persönlichkeit und den Zauber seiner Genialität; sein Wesen war wie der Klang einer Glocke, so hell in dem Glück bei der Arbeit und der Heiterkeit seines Gemütes, so voll in der Beständigkeit und Zuverlässigkeit, so rein in seinem idealen Streben und seiner Lebensauffassung.

Wie er gelebt hat, so starb er — als Philosoph. Wenige Stunden noch vor seinem Tode traf er die Anordnungen über die Korrektur seiner im Druck befindlichen Arbeit und überlegte, ob es sich empfehlen würde, seine unfertigen Manuskripte zu verwerten. Er sprach sein Bedauern über sein Schicksal aus, da er doch noch vieles hätte machen können; seiner letzten elektrodynamischen Arbeit aber würde es vielleicht zugute kommen, daß er zur Seite trete — man werde sie mehr lesen und mehr anerkennen. Zum Abschiednehmen verlangte er nach den Seinigen und nach mir.

Mehr als sechs Jahre hindurch haben wir, seine nächsten mathematischen Kollegen, jeden Donnerstag pünktlich drei Uhr mit ihm zusammen den mathematischen Spaziergang auf den Hainberg gemacht — auch den letzten Donnerstag vor seinem Tode, wo er uns mit besonderer Lebhaftigkeit von den neuen Fortschritten seiner elektrodynamischen Untersuchungen erzählte: den Donnerstag darauf — wiederum um drei Uhr — gaben wir ihm das letzte Geleit. Dienstag, den 12. Januar, mittags, war er einer Blinddarmentzündung erlegen; bei dem bösartigen Charakter, mit dem die Krankheit auftrat, hatte auch die Sonntag Nacht ausgeführte Operation nicht mehr helfen können.

Jäh hat ihn der Tod von unserer Seite gerissen. Was uns aber der Tod nicht nehmen kann, das ist sein edles Bild in unserem Herzen und das Bewußtsein, daß sein Geist in uns fortwirkt.

ZUR THEORIE
DER QUADRATISCHEN FORMEN

I.

Grundlagen für eine Theorie der quadratischen Formen mit ganzzahligen Koeffizienten.

(Mémoires présentés par divers savants à l'Académie des Sciences de l'Institut national de France, Tome XXIX, No. 2. 1884.)

(Vorbemerkung des Herausgebers. Diese Abhandlung ist von Minkowski im Mai 1882 der Pariser Académie des Sciences als Preisschrift, und zwar in Gestalt eines in deutscher Sprache verfaßten Manuskriptes, eingereicht worden (vgl. Comptes Rendus, Bd. 96 (1883), pp. 879 – 883). In den Mémoires des Savants étrangers ist sie, von Minkowski selbst ins Französische übertragen, unter dem Titel „Mémoire sur la théorie des formes quadratiques à coefficients entiers“ erschienen. Die vorliegende deutsche Ausgabe folgt überall dort, wo der gedruckte französische Text als unmittelbare Übersetzung des ursprünglichen deutschen Manuskriptes gelten kann, dem letzteren; an den zahlreichen Stellen, an denen beide voneinander abweichen, ist der französische Text als maßgebend angesehen und ins Deutsche rückübersetzt worden. Ein Verzeichnis der hauptsächlichsten dieser Stellen befindet sich am Schlusse der Abhandlung. Von dem Herausgeber herrührende Zusätze und Anmerkungen sind durch doppelte eckige Klammern [[]] kenntlich gemacht.)

[[Begleitschreiben an die Académie des Sciences.]]

J'ai l'honneur de soumettre au jugement de l'Académie mon Mémoire ci-joint, intitulé: «*Fondements pour une théorie des formes quadratiques à coefficients numériques*» comme ouvrage de concours au *Grand Prix des sciences mathématiques*, proposé pour la solution de la question: «Théorie de la décomposition des nombres entiers en une somme de cinq carrés.»

Tout en espérant avoir réussi à résoudre la question proposée, et à la généraliser en même temps, je dois demander l'indulgence de l'Académie en plus d'un rapport.

Arrivé il y a peu de jours seulement à ce point d'ample développement de ces attrayantes théories, j'ai dû y interrompre mon ouvrage, voyant qu'il ne me restait pour le faire arriver au terme du 1^{er} juin, que juste le temps de le mettre au net. C'est pourquoi avant tout il m'a été impossible d'en faire une traduction en français, ce qui, du reste

à l'état actuel de mes connaissances dans cette langue, je n'aurais su faire que d'une façon assez imparfaite.

Cette insuffisance de temps m'a empêché d'épuiser tout le matériel que j'ai trouvé pour cette question intéressante, ainsi que de donner à la rédaction tous ces soins que j'aurais voulu y porter, pour rendre mon ouvrage aussi en fait de forme, plus digne du noble forum, au jugement duquel j'ose le soumettre.

Je prie l'Académie de bien vouloir excuser ces faiblesses et de daigner examiner mon ouvrage tel que je suis forcé à le présenter.

L'auteur

du mémoire portant l'épigraphe: «Rien
n'est beau que le vrai, le vrai seul est
aimable.»

(Koenigsberg), le 29 mai 1882.

Durch die von der Académie des Sciences gestellte Aufgabe „Théorie de la décomposition des nombres entiers en une somme de cinq carrés“ angeregt, unternahm ich eine genauere Untersuchung der allgemeinen quadratischen Formen mit ganzzahligen Koeffizienten. Ich ging dabei von dem natürlichen Gedanken aus, daß die Zerlegung einer Zahl in eine Summe von fünf Quadraten in ähnlicher Weise von den quadratischen Formen mit vier Variablen abhängen würde, wie bekanntlich die Zerlegung einer Zahl in eine Summe von drei Quadraten von den quadratischen Formen mit zwei Variablen abhängt. Diese Untersuchung hat mir in der Tat die gewünschten Resultate über die Zerlegung einer Zahl in eine Summe von fünf Quadraten geliefert. Indessen erscheinen diese Resultate bei der großen Allgemeinheit der von mir gefundenen Sätze nicht überall als das eigentliche Hauptziel der vorliegenden Arbeit; sie stellen vielmehr nur ein Beispiel für die gewonnenen umfangreichen Theorien dar. Wenn daher viele der nachfolgenden Betrachtungen nicht immer unmittelbar auf das Thema der Preisfrage hinweisen, so wage ich dennoch zu hoffen, daß die Akademie nicht der Ansicht sein werde, ich würde mehr gegeben haben, wenn ich weniger gegeben hätte.

Ich teile kurz die bemerkenswertesten Sätze dieser Arbeit mit. — Es sei

$$f = \sum_{i,k=1}^n a_{ik} x_i x_k$$

eine quadratische Form mit ganzzahligen Koeffizienten a_{ik} und von einer nicht-verschwindenden Determinante Δ , welche sich durch eine reelle Transformation in eine Summe von $n - I$ positiven und I negativen Qua-

draten transformieren lasse. Die Zahl I heißt der Index der Form f . Das System

$$A = \begin{pmatrix} a_{11}, & a_{12}, & \dots, & a_{1n} \\ a_{21}, & a_{22}, & \dots, & a_{2n} \\ \cdot & \cdot & \cdot & \cdot \\ a_{n1}, & a_{n2}, & \dots, & a_{nn} \end{pmatrix}$$

nennen wir das quadratische System der Form f . Der größte Teiler aller h -reihigen Unterdeterminanten des Systemes A sei d_{h-1} , der größte Teiler aller doppelt genommenen unsymmetrischen und einfach genommenen symmetrischen h -reihigen Determinanten von A sei gleich $\sigma_h d_{h-1}$. σ_h ist gleich 1 oder gleich 2 ($\sigma_n = 1$, $d_{n-1} = (-1)^I \cdot \Delta$). Ist

$$g = \sum_{i,k=1}^n b_{ik} y_i y_k$$

eine zweite quadratische Form, welche durch eine ganzzahlige Substitution von der Determinante 1 aus f hervorgeht, so heißt g der Form f äquivalent, und es gelten, wenn der Form g die Zahlen c_{h-1} und ϱ_h [[in derselben Weise wie d_{h-1} und σ_h der Form f]] angehören, die Beziehungen

$$\sigma_h = \varrho_h, \quad d_{h-1} = c_{h-1}.$$

Wir betrachten hauptsächlich Formen f , für welche der größte Teiler der Koeffizienten a_{ik} , nämlich d_0 , gleich 1 ist, sogenannte primitive Formen. Für eine primitive Form setzen wir

$$\begin{aligned} d_1 &= o_1, \\ d_2 &= o_1^2 o_2, \\ &\cdot \quad \cdot \quad \cdot \quad \cdot \quad \cdot, \\ d_{n-1} &= o_1^{n-1} o_2^{n-2} \dots o_{n-1}. \end{aligned}$$

Die Zahlen o_h und σ_h nennen wir die Invarianten der (primitiven) Form f . Alle Formen f , für welche die $2(n-1) + 1$ Zahlen

$$\left(\begin{matrix} \sigma_1, \sigma_2, \dots, \sigma_{n-1} \\ o_1, o_2, \dots, o_{n-1} \end{matrix} \right), I$$

gleiche Werte haben, fassen wir in eine Ordnung zusammen. Zwei äquivalente Formen gehören derselben Ordnung an.

Die Existenz einer Ordnung ist (wenn wir $\sigma_0 = 1$, $o_0 = 0$; $\sigma_n = 1$, $o_n = 0$ setzen) an die folgenden Hauptbedingungen gebunden:

- I. Die Größen o_h sind ganze Zahlen.
- II. Die Zahlen σ_h und $\sigma_{h-1} \cdot o_h \cdot \sigma_{h+1}$ sind nicht gleichzeitig durch 2 teilbar (σ_h ist relativ prim zu $\sigma_{h-1} \cdot o_h \cdot \sigma_{h+1}$).
- III. Die Größen $\frac{\sigma_{h-1} o_h}{\sigma_{h+1}}$ und $\frac{\sigma_{h+1} o_h}{\sigma_{h-1}}$ sind ganze Zahlen.

Außer diesen Hauptbedingungen besteht im Falle, daß nicht sämtliche $n - 1$ Zahlen $\sigma_{h-1} o_h \sigma_{h+1}$ Quadratzahlen sind, die einzige Nebenbedingung:

Wenn $n \equiv 0 \pmod{2}$ und

$$\sigma_1 = 2, \sigma_2 = 1, \dots, \sigma_{n-2} = 1, \sigma_{n-1} = 2$$

wird, so ist, wenn wir die in den Zahlen o_h aufgehenden Potenzen von 2 gleich 2^{w_h} setzen,

$$\prod_{j=1}^{\frac{n}{2}} \frac{o_{2j-1}}{2^{w_{2j-1}}} \equiv (-1)^{I + \frac{n}{2}} \pmod{4}.$$

Im Falle, daß die sämtlichen $n - 1$ Zahlen $\sigma_{h-1} o_h \sigma_{h+1}$ Quadrate sind, treten noch mehrere einfache Nebenbedingungen hinzu. Alle Ordnungen, welche mit den von uns aufgestellten Bedingungen verträglich sind, enthalten wirklich Formenklassen.

Zum Beweise dieser Sätze in betreff der Formenordnungen führen wir den folgenden, sehr fruchtbaren und leicht auf Formen von beliebig hohem Grade auszudehnenden Begriff ein: Eine quadratische Form

$$R = \sum_{i,k=1}^n R_{ik} x_i x_k$$

heißt ein Rest einer Formenklasse f in bezug auf einen Modul N , wenn sich in dieser Formenklasse eine Form

$$\Phi = \sum_{i,k=1}^n A_{ik} x_i x_k$$

vorfindet, für welche die sämtlichen Kongruenzen

$$A_{ik} \equiv R_{ik} \pmod{N}$$

statthaben. — Wir können den Rest R der Formenklasse f so wählen, daß er nach dem Modul N in eine Summe von mehreren Einzelformen mit einer oder zwei Variablen zerfällt, und wir nennen einen Rest R , für welchen die Anzahl dieser Einzelformen den größtmöglichen Wert erhält und in welchem die n Variablen in bestimmter Weise geordnet sind, einen Hauptrest der Klasse f . Ein Hauptrest R einer Klasse f hängt, wie wir zeigen, nur von den Resten seiner n mittleren*) Koeffizienten ab.

Wir bezeichnen die aus den ersten h Reihen einer Form f gebildeten symmetrischen Unterdeterminanten durch $\sigma_h d_{h-1} f_h$. Wir können in jeder Formenklasse f eine Form φ bestimmen, für welche die Zahlen φ_h zu den Zahlen $2 o_1 o_2 \dots o_{n-1} \cdot \varphi_{h-1} \varphi_{h+1}$ relativ prim ausfallen. Eine derartige Form heißt eine charakteristische Form der Klasse f . Die aus den ersten

*) [[Die R_{ii} werden als mittlere, die R_{ik} ($i \neq k$) als seitliche Koeffizienten einer quadratischen Form $R = \sum R_{ik} x_i x_k$ bezeichnet.]]

h [[Horizontal- und Vertikal-]] Reihen des quadratischen Systems von φ gebildete Form $\{\varphi_h\}$ von h Variablen möge den Index I_h besitzen. Die aus den $n+1$ Zahlen

$$I_0 = 0; \quad I_1, \quad I_2, \dots, \quad I_{n-1}; \quad I_n = I$$

entstehenden Einheiten $\varepsilon_h = (-1)^{I_h}$ stellen die Vorzeichen der Zahlen φ_h vor, und es sind mithin die Größen $\varepsilon_h \varphi_h$ positiv.

Zwei Formenklassen f und g , welche irgendeinen gleichen Formenrest für einen Modul N besitzen, enthalten lauter gleiche Formenreste für den Modul N . Wir nennen zwei derartige Formenklassen kongruent für den Modul N . Alle Formenklassen, welche für einen jeden Modul N kongruent sind und welche einen gleichen Index I besitzen, fassen wir in ein Formengenus zusammen. Äquivalente Formen gehören demselben Genus an.

Zwei Formenklassen, welche in einem und demselben Genus enthalten sind, besitzen dieselbe Ordnung.

Für alle charakteristischen Formen φ der verschiedenen, in einem Genus G enthaltenen Formenklassen f besitzen,

$$\left\{ \begin{array}{l} \text{I. wenn } \sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{p} \text{ ist und } p \text{ eine ungerade Primzahl bedeutet, die Einheiten} \\ \quad \left(\frac{\varphi_h}{p} \right), \\ \text{II. wenn } \sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{4} \text{ ist, die drei Einheiten} \\ \quad (-1)^{\frac{\varphi_h-1}{2}}, \left(\frac{\varphi_{h-1}}{\varepsilon_h \varphi_h} \right) \cdot (-1)^{\frac{I_h(I_h-1)}{1 \cdot 2}}, \left(\frac{\varphi_{h+1}}{\varepsilon_h \varphi_h} \right) \cdot (-1)^{\frac{I_h(I_h+1)}{1 \cdot 2}}; \\ \text{III. wenn } \sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{8} \text{ ist, die Einheiten} \\ \quad \left(\frac{2}{\varphi_h} \right) \end{array} \right. \quad (\mathfrak{G})$$

die nämlichen Werte.

Umgekehrt gehören zwei Formen f wirklich stets einem und demselben Genus an, sobald für die charakteristischen Formen φ dieser Formenklassen die vorstehenden Einheiten die nämlichen Werte besitzen.

Zwischen den Einheiten $(\mathfrak{G})$, die wir die Charaktere des Genus G nennen, bestehen alle diejenigen Bedingungen, welche sich aus den Kongruenzen

$$k) \quad -\sigma_{h-1} o_h \sigma_{h+1} \cdot \varphi_{h-1} \varphi_{h+1} \equiv X_h^2 \pmod{\sigma_h^2 \varphi_h}$$

erschließen lassen, also insbesondere die Bedingungen

$$\left\{ \left(\frac{\varphi_{h-1}}{\varepsilon_h \varphi_h} \right) \cdot (-1)^{\frac{I_h(I_h-1)}{2}} \right\} \cdot \left\{ \left(\frac{\varphi_{h+1}}{\varepsilon_h \varphi_h} \right) \cdot (-1)^{\frac{I_h(I_h+1)}{2}} \right\} \\ = \left(\frac{\sigma_{h-1} 2^{\omega_h} \sigma_{h+1}}{\varphi_h} \right) \cdot (-1)^{\frac{\varphi_{h-1}}{2} \cdot \frac{e_h+1}{2}} \cdot \left(\frac{\varphi_h}{e_h} \right)$$

$$[o_h = 2^{\omega_h} \cdot e_h, \quad e_h \equiv 1 \pmod{2}]$$

und

$$\prod_{h=1}^{n-1} \left(\frac{\sigma_{h-1} 2^{\omega_h} \sigma_{h+1}}{\varphi_h} \right) \cdot (-1)^{\sum_{h=1}^{n-1} \frac{\varphi_h - 1}{2} \cdot \frac{e_h + 1}{2}} \prod_{h=1}^{n-1} \left(\frac{\varphi_h}{e_h} \right) \\ = (-1)^{\left[\frac{f}{2} \right]} \cdot (-1)^{\frac{\varphi_0 - 1}{2} \cdot \frac{\varphi_1 - 1}{2} + \frac{\varphi_1 - 1}{2} \cdot \frac{\varphi_2 - 1}{2} + \dots + \frac{\varphi_{n-1} - 1}{2} \cdot \frac{\varphi_n - 1}{2}}.$$

Wenn die Charaktere eines Genus G allen Bedingungen genügen, welche sich aus den Kongruenzen k) ergeben, so besitzt das Genus G wirklich Formen. (Insbesondere folgt aus diesem Umstand die bis jetzt noch nicht bewiesene Tatsache, daß die definiten positiven Formen von der Determinante 1 und einer Variablenzahl größer oder gleich 8 mehr als eine Formenklasse enthalten.)

Es sei f eine primitive Form. Wir bezeichnen die $(n-1)$ -reihigen Minoren der Determinante Δ in bekannter Weise durch

$$\frac{\partial \Delta}{\partial a_{ik}} = A_{ik},$$

und wir setzen

$$A_{ik} = (-1)^i d_{n-2} \cdot a'_{n-i+1, n-k+1}.$$

Die Form

$$f' = \sum_{i,k=1}^n a'_{ik} x'_i x'_k$$

nennen wir dann die der Form f adjungierte Form.

Ist eine Form f' einer Form f adjungiert, so ist auch der Form f' die Form f adjungiert.

Besitzt die Form f eine Ordnung

$$\begin{pmatrix} \sigma_1, \sigma_2, \dots, \sigma_{n-1} \\ o_1, o_2, \dots, o_{n-1} \end{pmatrix}, I,$$

so besitzt ihre adjungierte Form f' eine Ordnung

$$\begin{pmatrix} \sigma'_1, \sigma'_2, \dots, \sigma'_{n-1} \\ o'_1, o'_2, \dots, o'_{n-1} \end{pmatrix}, I',$$

für welche

$$\sigma'_h = \sigma_{n-h}, \quad o'_h = o_{n-h}; \quad I' = I$$

wird. — Der charakteristischen Form φ' der Klasse f' ist eine charakteristische Form φ der Klasse f adjungiert, für welche die Beziehungen

$$\varphi'_h = \varphi_{n-h} \cdot (-1)^i \quad \text{und} \quad I'_h + I_{n-h} = I$$

bestehen. Infolge dieser Beziehungen können die Charaktere und demgemäß die Genera adjungierter Formen sofort auseinander hergeleitet werden.

Die Beweise der soeben entwickelten Sätze bilden den Inhalt des ersten Teiles der vorliegenden Arbeit, in welchem alle quadratischen

Formen f von nicht-verschwindender Determinante ohne Beschränkung behandelt sind.

In dem zweiten Teile dieser Arbeit betrachte ich die Darstellung von Formen mit niedrigerer Variablenzahl durch Formen mit höherer Variablenzahl und insbesondere die Darstellung von ganzen Zahlen durch Formen. Ich verallgemeinere und erweitere die von Gauß über Darstellungen von Zahlen und binären Formen durch ternäre Formen gegebenen Sätze und betrachte im weiteren Verlaufe der Untersuchung besonders definite positive Formen. Zuletzt wende ich mich zu einer Bestimmung des Maßes einiger besonders ausgezeichneter definiter Genera und gewinne dadurch die verlangten Sätze über die Darstellung einer Zahl durch eine Summe von fünf Quadraten. Schließlich füge ich noch ohne Beweis ein sehr allgemeines Theorem in betreff des Maßes eines definiten positiven Genus von n Variablen hinzu, aus welchem sich sämtliche von Eisenstein in diesem Gebiete aufgestellten Sätze als spezielle Fälle ergeben. — Infolge der Nähe des von der Akademie gegebenen Termines mußte ich mich in dem zweiten Teile an manchen Stellen mit kurzen Andeutungen begnügen, und es sind dadurch einige Mängel in der Redaktion desselben veranlaßt, welche ich zu entschuldigen bitte.

Erster Teil.

Über die Reste quadratischer Formen.

Kap. I. Klassen quadratischer Formen. — Index, Invarianten und Ordnung einer Form.

Wir betrachten quadratische Formen $f = \sum_{i,k=1}^n a_{ik} x_i x_k = \{a_{ik}\}$, welche ganzzahlige Koeffizienten und eine nicht-verschwindende Determinante besitzen.

Die Anwendung einer beliebigen ganzzahligen Substitution

$$S: \quad x_i = \sum_l s_{il} y_l \quad (i, l = 1, 2, \dots, n)$$

auf die quadratische Form f führt zu einer neuen Form

$$g = \bar{S} f S = \sum_{l,m=1}^n b_{lm} y_l y_m. *)$$

Die Form g heißt in der Form f *enthalten*.

*) Wir bezeichnen mit $\bar{S}$ dasjenige System, das sich aus S durch Vertauschung der Horizontal- und Vertikalreihen ableitet.

Wir beschäftigen uns insbesondere mit Substitutionen S , deren Determinante $= 1$ ist. In diesem Falle kann die Form g durch eine zweite ganzzahlige Substitution von der Determinante 1 in f zurückgeführt werden. Für $|S| = 1$ ist daher sowohl g in f als f in g enthalten. Mit Rücksicht auf diese ihre gegenseitige Beziehung nennen wir alsdann f und g äquivalent. Um die Äquivalenz zweier Formen auszudrücken, bedienen wir uns mitunter des Zeichens $\sim$ ($f \sim g$). Es besteht der Fundamentalsatz:

A. Wenn zwei Formen einer dritten äquivalent sind, so sind sie untereinander äquivalent.

Denn ist

$$g_1 = \bar{S}_1 f S_1, \quad g_2 = \bar{S}_2 f S_2,$$

so haben wir

$$g_1 = (\bar{S}_2^{-1} \bar{S}_1) \cdot g_2 \cdot (S_2^{-1} S_1) \sim g_2.$$

Die Gesamtheit aller einer bestimmten Form f äquivalenten Formen nennen wir eine *Formenklasse*. Infolge des Satzes A. sind zwei Formen, welche derselben Klasse angehören, untereinander äquivalent, während zwei Formen aus verschiedenen Klassen nicht äquivalent sind. Wir werden eine gegebene Form selten als besonderes, durch seine $\frac{n(n+1)}{2}$ Koeffizienten bestimmtes Individuum betrachten; meistens wird sie uns nur als ein *Repräsentant* der ihr entsprechenden Formenklasse gelten.

Eine jede Form f von nicht-verschwindender Determinante kann, wie man weiß, durch eine lineare Transformation mit reellen (keineswegs immer ganzzahligen) Koeffizienten in eine Summe von n zum Teil positiven, zum Teil negativen Quadraten übergeführt werden. Es mögen bei irgendeiner solchen Transformation $I = I(f)$ Quadrate das negative und $n - I$ Quadrate das positive Vorzeichen erhalten; alsdann ist nach einem bekannten Satze die Zahl I für die Form f charakteristisch, und es wird eine jede Darstellung von f durch eine Summe von n positiven oder negativen Quadraten die Gestalt

$$f = - \sum_{h=1}^I x_h^2 + \sum_{h=1}^{n-I} \bar{x}_h^2$$

erhalten. Wir gewinnen hieraus leicht für die algebraische Äquivalenz zweier Formen f und g die Bedingung

$$I(f) = I(g).$$

Die Zahl $I(f)$ heißt der *Index* der Form f .

Die Form $g = \{b_{im}\}$ sei in der Form $f = \{a_{ik}\}$ mittels einer Substitution $S = \{s_i^j\}$ enthalten. Wir bezeichnen die Subdeterminanten h^{ten} Grades

$$\begin{vmatrix} a_{i_1 k_1} & a_{i_1 k_2} & \dots & a_{i_1 k_h} \\ a_{i_2 k_1} & a_{i_2 k_2} & \dots & a_{i_2 k_h} \\ \dots & \dots & \dots & \dots \\ a_{i_h k_1} & a_{i_h k_2} & \dots & a_{i_h k_h} \end{vmatrix}, \begin{vmatrix} b_{i_1 m_1} & b_{i_1 m_2} & \dots & b_{i_1 m_h} \\ b_{i_2 m_1} & b_{i_2 m_2} & \dots & b_{i_2 m_h} \\ \dots & \dots & \dots & \dots \\ b_{i_h m_1} & b_{i_h m_2} & \dots & b_{i_h m_h} \end{vmatrix}, \begin{vmatrix} s_{i_1}^{l_1} & s_{i_1}^{l_2} & \dots & s_{i_1}^{l_h} \\ s_{i_2}^{l_1} & s_{i_2}^{l_2} & \dots & s_{i_2}^{l_h} \\ \dots & \dots & \dots & \dots \\ s_{i_h}^{l_1} & s_{i_h}^{l_2} & \dots & s_{i_h}^{l_h} \end{vmatrix}$$

durch

$$A \begin{pmatrix} i_1, i_2, \dots, i_h \\ k_1, k_2, \dots, k_h \end{pmatrix}, B \begin{pmatrix} l_1, l_2, \dots, l_h \\ m_1, m_2, \dots, m_h \end{pmatrix}, S \begin{pmatrix} l_1, l_2, \dots, l_h \\ i_1, i_2, \dots, i_h \end{pmatrix},$$

oder, wofern die besondere Wahl der h Horizontal- und Vertikalreihen unter den gegebenen n Horizontal- und Vertikalreihen gleichgültig ist, durch

$$A_h, B_h, S_h \text{ oder } S_h^0.$$

Wir gebrauchen für die symmetrischen A_h und B_h die Buchstaben F_h und G_h und für die unsymmetrischen A_h und B_h die Buchstaben P_h und Q_h .

Nach einem bekannten Determinantensatz ist

$$B \begin{pmatrix} l_1, l_2, \dots, l_h \\ m_1, m_2, \dots, m_h \end{pmatrix} = \frac{\sum_{i,k}^{1,n} A \begin{pmatrix} i_1, i_2, \dots, i_h \\ k_1, k_2, \dots, k_h \end{pmatrix} S \begin{pmatrix} l_1, l_2, \dots, l_h \\ i_1, i_2, \dots, i_h \end{pmatrix} S \begin{pmatrix} m_1, m_2, \dots, m_h \\ k_1, k_2, \dots, k_h \end{pmatrix}}{(1 \cdot 2 \dots h) \cdot (1 \cdot 2 \dots h)}.$$

In der Entwicklung eines G_h sehen wir die F_h mit einem Faktor S_h^2 , die P_h mit einem Faktor $2S_h S_h^0$ behaftet; in der Entwicklung eines Q_h weisen sowohl die F_h als die P_h einen Faktor $S_h S_h^0$ auf:

$$G_h = \sum F_h S_h^2 + \sum P_h 2S_h S_h^0, \\ Q_h = \sum F_h S_h S_h^0 + \sum P_h S_h S_h^0.$$

Wir schließen aus dieser Bemerkung, daß der größte Divisor der F_h , P_h in dem größten Divisor der G_h , Q_h und der größte Divisor der F_h , $2P_h$ in dem größten Divisor der G_h , $2Q_h$ aufgeht.

Den größten positiven Divisor der F_h , P_h setzen wir gleich d_{h-1} und den größten positiven Divisor der F_h , $2P_h$ gleich $\sigma_h d_{h-1}$. Die Zahl σ_h stellt uns dann den größten Divisor der Größen $\frac{F_h}{d_{h-1}}$, $2 \frac{P_h}{d_{h-1}}$ vor. Da der größte Teiler der $\frac{F_h}{d_{h-1}}$, $\frac{P_h}{d_{h-1}}$ die Einheit ist, so nimmt σ_h den Wert 1 an, wenn von den Zahlen $\frac{F_h}{d_{h-1}}$ wenigstens eine ungerade ist, und σ_h wird gleich 2, wenn alle $\frac{F_h}{d_{h-1}}$ gerade ausfallen. Die Zahl d_{n-1} ist dem absoluten Wert der Determinante $|a_{ik}|$ gleich, und es wird daher

$$|a_{ik}| = (-1)^t \cdot d_{n-1};$$

ferner ist $\sigma_n = 1$.

Enthält nicht nur die Form f die Form g , sondern auch die Form g die Form f , sind also f und g äquivalent, so geht einerseits der größte positive Divisor $d_{h-1}(f)$ der F_h, P_h in dem größten positiven Divisor $d_{h-1}(g)$ der G_h, Q_h und der größte positive Divisor $\sigma_h(f)d_{h-1}(f)$ der $F_h, 2P_h$ in dem größten positiven Divisor $\sigma_h(g)d_{h-1}(g)$ der $G_h, 2Q_h$ auf; andererseits geht auch $d_{h-1}(g)$ in $d_{h-1}(f)$ und $\sigma_h(g)d_{h-1}(g)$ in $\sigma_h(f)d_{h-1}(f)$ auf. Für äquivalente Formen f und g bestehen demnach die Beziehungen

$$(1) \quad \sigma_h(f) = \sigma_h(g); \quad d_{h-1}(f) = d_{h-1}(g).$$

Wir nennen eine Form $f = \{a_{ik}\}$ *primitiv in bezug auf einen Modul* N , sobald die Koeffizienten a_{ik} einen größten gemeinsamen Teiler besitzen, welcher zu N relativ prim ist. Eine beliebige Form ist nach dieser Definition primitiv in bezug auf einen Modul N , sobald sie primitiv in bezug auf alle in N enthaltenen Primfaktoren ist.

In dem Folgenden werden wir fast ausschließlich Formen f betrachten, für welche der größte Teiler der a_{ik} , nämlich d_0 , gleich 1 wird und die wir kurz *primitive* Formen heißen. Eine nicht-primitive Form $f = \{a_{ik}\}$ ist dem Produkt aus der Zahl d_0 und der primitiven Form $\frac{f}{d_0} = \left\{ \frac{a_{ik}}{d_0} \right\}$ gleich.

Jeder primitiven Form f entsprechen $2(n-1)$ Zahlen

$$\begin{aligned} \sigma_1, \sigma_2, \dots, \sigma_{n-2}, \sigma_{n-1}, \\ d_1, d_2, \dots, d_{n-2}, d_{n-1}. \end{aligned}$$

Wir führen ferner $n-1$ Größen o_h durch die Gleichungen

$$\begin{aligned} d_1 &= o_1, \\ d_2 &= o_1^2 o_2, \\ &\dots \dots \dots \\ d_{n-1} &= o_1^{n-1} o_2^{n-2} \dots o_{n-1}, \end{aligned} \quad o_h = \frac{d_h d_{h-2}}{d_{h-1}^2}$$

ein.

Gemäß den Formeln (1) sind alle o_h und σ_h *Invarianten* der Klasse f . Wir unterscheiden die o_h und σ_h als Invarianten $o(f)$ und Invarianten $\sigma(f)$.

Alle Formenklassen, welche die Größen σ_h, o_h gemein haben und für welche der Index $I(f)$ denselben Wert I annimmt, fassen wir in eine *Ordnung*

$$\left(\begin{array}{c} \sigma_1, \sigma_2, \dots, \sigma_{n-2}, \sigma_{n-1} \\ o_1, o_2, \dots, o_{n-2}, o_{n-1} \end{array} \right), \quad I$$

zusammen.

Unsere nächste Aufgabe wird in der Aufsuchung der für die Existenz einer Ordnung erforderlichen Bedingungen bestehen.

Kap. II. Formenreste. — Die Invarianten $o(f)$ sind ganze Zahlen.

Bei allen Fragen, in welchen es sich weniger um die numerischen Werte der Unterdeterminanten einer Form $f = \{a_{ik}\}$ als um die Teilbarkeit dieser Unterdeterminanten durch gegebene Zahlen N handelt, kommen für uns offenbar nur die Reste der $\frac{n(n+1)}{2}$ Koeffizienten a_{ik} nach den Moduln N in Betracht, und wir können daher diese Koeffizienten durch beliebige, ihnen nach den Moduln N kongruente Zahlen ersetzen.

Eine Form $R = \{R_{ik}\}$ heißt ein Rest der Formenklasse f in bezug auf den Modul N , sobald in dieser Klasse eine Form $\Phi = \{A_{ik}\}$ anzutreffen ist, für welche sämtliche Kongruenzen

$$A_{ik} \equiv R_{ik} \pmod{N}$$

erfüllt sind. Wir gebrauchen alsdann die Bezeichnung

$$\Phi \equiv \{R_{ik}\} \pmod{N}.$$

Wir werden für eine jede Formenklasse eine Reihe besonders ausgezeichnete Formenreste bestimmen, zu denen man gelangt, wenn man auf eine beliebige Form der betreffenden Klasse Substitutionen von der folgenden Art ausübt:

$$\begin{array}{lll} S_{(i,k)}^{\pm 1}: & x_i = x'_i, \quad x_k = \pm x'_i + x'_k, & x_h = x'_h \quad (h \neq i, k) \\ \text{und} & & \\ P_{(l,m)}: & x_i = x'_m, \quad x_m = -x'_i, & x_h = x'_h \quad (h \neq l, m) \\ \text{und} & & \\ & x_i = x'_i + \sum_{(k>i)} s_k x'_k, & x_h = x'_h \quad (h \neq i). \end{array}$$

Zunächst verschafft uns die Betrachtung der so gewonnenen Klassenreste den Satz:

B. Für eine jede Ordnung, welche wirklich Formen enthält, werden die Invarianten o_h ganze Zahlen.

Wir können diesen Satz auch in der folgenden veränderten Weise aussprechen:

C. Ist die quadratische Form $f = \{a_{ik}\}$ in bezug auf eine Primzahl q primitiv und sind die höchsten, in den Zahlen $d_1, d_2, \dots, d_{n-1}$ der Form f aufgehenden Potenzen von q resp. $q^{\partial_1}, q^{\partial_2}, \dots, q^{\partial_{n-1}}$, so fallen die Größen $\omega_h = (\partial_h - \partial_{h-1}) - (\partial_{h-1} - \partial_{h-2})$ nicht negativ aus.

Die ω_h sind mit den ∂_h durch die Gleichungen

$$\partial_h = h\omega_1 + (h-1)\omega_2 + \dots + \omega_h$$

verbunden. Die Potenzen q^{ω_h} sind offenbar diejenigen Potenzen von q , welche in den Größen $o_h = \frac{d_h}{d_{h-1}} : \frac{d_{h-1}}{d_{h-2}}$ aufgehen; sobald also keine der

Zahlen ω_h negativ ist, müssen die o_h sämtlich ganze Zahlen sein. Die erste von 0 verschiedene der Zahlen ω_h ist gewiß positiv, da sie gleich der ersten von 0 verschiedenen Zahl ∂_h ist. — Wir werden die Gültigkeit des Satzes C. für Formen von n Variablen aus der Annahme der Gültigkeit desselben für Formen von weniger als n Variablen herleiten. Hierdurch ist dann zugleich seine Allgemeingültigkeit dargetan; denn da dieser Satz für $n = 1$ gewiß richtig ist — in diesem Fall existiert überhaupt kein o_h —, wird er sofort auch für $n = 2, 3, \dots, n$ erwiesen sein.

Für den Beweis des Satzes C. genügt die Einführung besonderer Formenreste für Moduln, welche Potenzen der Primzahl q sind. Nur müssen wir den Fall einer ungeraden Primzahlpotenz $q^t = p^t$ von dem Falle einer Potenz $q^t = 2^t$ sondern.*)

I. Wir beginnen mit dem Falle einer ungeraden Primzahl $q = p$.

1. Jede in bezug auf p primitive Form $f = \sum_{i,k=1}^n a_{ik} x_i x_k$ von n Variablen ist einer Form

$$f_{(1)} \equiv \begin{pmatrix} \alpha, & 0, & \dots, & 0 \\ 0, & p^{\omega_1} a_{11}^{(1)}, & \dots, & p^{\omega_1} a_{1,n-1}^{(1)} \\ \cdot & \cdot & \cdot & \cdot \\ 0, & p^{\omega_1} a_{n-1,1}^{(1)}, & \dots, & p^{\omega_1} a_{n-1,n-1}^{(1)} \end{pmatrix} \pmod{p^f},$$

$$f_{(1)} \equiv \alpha \xi^2 + p^{\omega_1} \sum_{i,k=1}^{n-1} a_{ik}^{(1)} x_i^{(1)} x_k^{(1)} \pmod{p^f}$$

äquivalent, deren erster Koeffizient α zu p relativ prim ist, während

$f^{(1)} = \sum_{i,k=1}^{n-1} a_{ik}^{(1)} x_i^{(1)} x_k^{(1)}$ eine in bezug auf p primitive Form von $n-1$ Variablen darstellt.

Beweis: Wenn f in bezug auf p primitiv ist, so sind die Zahlen a_{ii} , $2a_{ik}$ gewiß nicht alle durch p teilbar. Ebenso können auch die Zahlen a_{ii} , $a_{ii} \pm 2a_{ik} + a_{kk}$ nicht alle durch p teilbar sein; denn sie haben offenbar denselben größten gemeinsamen Teiler wie die a_{ii} , $2a_{ik}$. Sollten also die Koeffizienten a_{ii} der Form f sämtlich durch p teilbar sein, so würde eine der Zahlen $a_{ii} \pm 2a_{ik} + a_{kk}$ zu p relativ prim sein, und wir dürften auf f nur eine Substitution $S_{(i,k)}^{\pm 1}$ anwenden, um in der veränderten Form die zu p prime Zahl $a_{ii} \pm 2a_{ik} + a_{kk}$ als einen mittleren Koeffizienten zu erzielen. Wir können demnach von f voraussetzen, daß es einen durch p nicht teilbaren Koeffizienten a_{ii} besitzt. Durch Ausübung

*) Wir bezeichnen stets durch q eine beliebige und durch p eine ungerade Primzahl.

einer Substitution $P_{(1,t)}$ verwandelt sich f alsdann in einen Repräsentanten

$$\hat{f}_{(1)} \equiv \begin{pmatrix} \alpha, & E_1, & \dots, & E_{n-1} \\ E_1, & c_{11}, & \dots, & c_{1,n-1} \\ \dots & \dots & \dots & \dots \\ E_{n-1}, & c_{n-1,1}, & \dots, & c_{n-1,n-1} \end{pmatrix} \pmod{p^t},$$

in welchem α zu p prim ist. Auf $\hat{f}_{(1)}$ wenden wir die Substitution

$$S^{(1)} = \begin{pmatrix} 1, & K_1, & K_2, & \dots, & K_{n-1} \\ & 1, & 0, & \dots, & 0 \\ & & \dots & \dots & \dots \\ & & & & 1 \end{pmatrix}$$

an. Eine solche Substitution läßt α ungeändert, während sie E_h in

$$E_h + \alpha K_h$$

verwandelt. Bestimmen wir daher die K_h aus den gewiß lösbaren Kongruenzen

$$E_h + \alpha K_h \equiv 0 \pmod{p^t},$$

so geht die Form $\hat{f}_{(1)}$ durch die Substitution $S^{(1)}$ in eine Form

$$f_{(1)} \equiv \begin{pmatrix} \alpha, & 0, & \dots, & 0 \\ 0, & r_{11}^{(1)}, & \dots, & r_{1,n-1}^{(1)} \\ \dots & \dots & \dots & \dots \\ 0, & r_{n-1,1}^{(1)}, & \dots, & r_{n-1,n-1}^{(1)} \end{pmatrix} \pmod{p^t}$$

über. Ist $t > \omega_1$, so wird p^{ω_1} die größte Potenz von p sein, welche in allen $r_{ik}^{(1)}$ enthalten ist, und wenn wir $r_{ik}^{(1)} = p^{\omega_1} a_{ik}^{(1)}$ setzen, so haben wir die angegebene Gestalt von $f_{(1)}$ erlangt.

2. Der Form $f(\sim f_{(1)})$ teilen wir $n-1$ Zahlen $d_1, d_2, \dots, d_{n-1}$ zu; $n-2$ Zahlen $d_1^{(1)}, d_2^{(1)}, \dots, d_{n-2}^{(1)}$ von ähnlicher Bedeutung gehören der Form $f^{(1)} = \{a_{ik}^{(1)}\}$ an. Es bezeichnet $d_{h-1}^{(1)}$ für $f^{(1)}$ den größten gemeinsamen Teiler aller $A_h^{(1)}$. Es seien die höchsten in $d_1^{(1)}, d_2^{(1)}, \dots, d_{n-2}^{(1)}$ aufgehenden Potenzen von p resp. $p^{\delta_1^{(1)}}, p^{\delta_2^{(1)}}, \dots, p^{\delta_{n-2}^{(1)}}$. Für die Form $f_{(1)}$ gelten die Kongruenzen

$$A_{(1)h} \equiv \alpha \cdot p^{(h-1)\omega_1} A_{h-1}^{(1)}, \quad p^{h\omega_1} A_h^{(1)}, \quad 0 \pmod{p^t}.$$

Wir wollen den Modul p^t größer gewählt denken als die größte der Zahlen $p^{\delta_1}, p^{\delta_2}, \dots, p^{\delta_{n-1}}$. Dann muß die höchste Potenz von p , welche in allen $A_{(1)h}$ aufgeht, d. i. die Potenz $p^{\delta_{h-1}}$, zugleich die höchste Potenz von p sein, welche in allen $\alpha \cdot p^{(h-1)\omega_1} A_{h-1}^{(1)}, p^{h\omega_1} A_h^{(1)}$ aufgeht. Die höchste Potenz von p , welche in den $\alpha \cdot p^{(h-1)\omega_1} A_{h-1}^{(1)}$ aufgeht, ist offenbar $p^{(h-1)\omega_1 + \delta_{h-1}^{(1)}}$; die höchste Potenz von p , welche in den $p^{h\omega_1} A_h^{(1)}$ aufgeht,

ist $p^{h\omega_1 + \partial_{h-1}^{(1)}}$. Demnach wird ∂_{h-1} mit der kleineren der Zahlen

$$(h-1)\omega_1 + \partial_{h-2}^{(1)}, \quad h\omega_1 + \partial_{h-1}^{(1)}$$

übereinstimmen müssen.

Setzen wir jetzt die Gültigkeit unseres Satzes C. für die Form $f^{(1)}$ von $n-1$ Variablen voraus. Dann werden die durch die Gleichungen

$$\partial_h^{(1)} = h\omega_1^{(1)} + (h-1)\omega_2^{(1)} + \dots + \omega_h^{(1)}$$

bestimmten Größen $\omega_h^{(1)}$ positiv oder gleich Null ausfallen, und es wird

$$\partial_{h-2}^{(1)} \leq \partial_{h-1}^{(1)},$$

woraus sich wegen $\omega_1 \geq 0$

$$(h-1)\omega_1 + \partial_{h-2}^{(1)} \leq h\omega_1 + \partial_{h-1}^{(1)}$$

ergibt. Von diesen beiden Zahlen ist also jedenfalls die erste nicht größer als die zweite, und wir erhalten folglich

$$\partial_{h-1} = (h-1)\omega_1 + \partial_{h-2}^{(1)},$$

oder, wenn wir

$$\omega_{h-1}^{(1)} = \omega_h \quad (h > 1)$$

setzen und die Gleichungen ausführlicher schreiben:

$$\partial_h = h\omega_1 + (h-1)\omega_2 + \dots + \omega_h, \quad \omega_h \geq 0.$$

II. Wenn die Primzahl $q = 2$ ist, unterscheiden wir die Fälle $\sigma_1 = 1$ und $\sigma_1 = 2$.

$\sigma_1 = 1$. — 1. Jede in bezug auf 2 primitive Form $f = \sum_{i,k=1}^n a_{ik} x_i x_k$,

welche eine Invariante $\sigma_1 = 1$ besitzt, ist einer Form

$$f_{(1)} \equiv \begin{pmatrix} \alpha, & 0, & \dots, & 0 \\ 0, & 2^{\omega_1} a_{11}^{(1)}, & \dots, & 2^{\omega_1} a_{1,n-1}^{(1)} \\ \cdot & \cdot & \cdot & \cdot \\ 0, & 2^{\omega_1} a_{n-1,1}^{(1)}, & \dots, & 2^{\omega_1} a_{n-1,n-1}^{(1)} \end{pmatrix} \pmod{2^f},$$

$$f_{(1)} \equiv \alpha \xi^2 + 2^{\omega_1} \sum_{i,k=1}^{n-1} a_{ik}^{(1)} x_i^{(1)} x_k^{(1)} \pmod{2^f}$$

äquivalent, in welcher der erste Koeffizient α ungerade ist, während

$f^{(1)} = \sum_{i,k=1}^{n-1} a_{ik}^{(1)} x_i^{(1)} x_k^{(1)}$ eine in bezug auf 2 primitive Form vorstellt, welche

im Falle $\omega_1 = 0$ eine erste Invariante σ gleich 1 besitzt.

Beweis: Es befinden sich, wenn $\sigma_1 = 1$ ist, unter den Koeffizienten a_{ii} ungerade Größen. Außerdem muß es im Falle $\omega_1 = 0$ eine ungerade Determinante $a_{hk} a_{hh} - a_{hh}^2$ geben. Denn wären alle diese Determinanten gerade, so würden die Kongruenzen

und

$$a_{kk} a_{hh} \equiv a_{kh}^2 \pmod{2}$$

$$a_{kh}^2 a_{k_0 h_0}^2 \equiv a_{kk} a_{hh} a_{k_0 k_0} a_{h_0 h_0} \equiv a_{kh_0}^2 a_{k_0 h}^2 \pmod{2}$$

gelten, und folglich wären auch alle Determinanten $a_{kh} a_{k_0 h_0} - a_{kh_0} a_{k_0 h}$ gerade, so daß die höchste Potenz von 2, welche in allen diesen Determinanten enthalten ist, nämlich 2^{ω_i} , nicht gleich 1 sein könnte. Wir können im Falle $\omega_1 = 0$ ferner voraussetzen, daß die Form f für denselben Index i sowohl ein ungerades a_{ii} als auch eine ungerade Determinante $a_{ii} a_{hh} - a_{ih}^2$ besitzt. Denn wären in f die Indizes i der ungeraden a_{ii} alle von den Indizes k, h der ungeraden Zahlen $a_{kk} a_{hh} - a_{kh}^2$ verschieden, so hätte man gleichzeitig

$$\begin{aligned} a_{ii} &\equiv 1, & a_{kk} a_{hh} - a_{kh}^2 &\equiv 1; \\ a_{ii} a_{kk} - a_{ik}^2 &\equiv 0, & a_{ii} a_{hh} - a_{ih}^2 &\equiv 0, & a_{kk} &\equiv 0, & a_{hh} &\equiv 0, \end{aligned} \pmod{2},$$

d. h.

$$\begin{pmatrix} a_{ii} & a_{ik} & a_{ih} \\ a_{ki} & a_{kk} & a_{kh} \\ a_{hi} & a_{hk} & a_{hh} \end{pmatrix} \equiv \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \\ 0 & 1 & 0 \end{pmatrix} \pmod{2}.$$

Durch Ausübung einer Substitution $S_{(i,k)}^{\pm 1}$ auf f erhalte man eine Form, in welcher die Zahlen a_{ii} und $a_{ii} a_{hh} - a_{ih}^2$ bzw. durch die Zahlen

$$a_{ii} \pm 2a_{ik} + a_{kk}$$

und

$$(a_{ii} a_{hh} - a_{ih}^2) \pm 2(a_{ik} a_{hh} - a_{ih} a_{kh}) + (a_{kk} a_{hh} - a_{kh}^2)$$

ersetzt erscheinen, welche beide ungerade sind. Vermittels einer Substitution $P_{(1,i)}$ können wir jetzt f in eine Form

$$\widehat{f}_{(1)} \equiv \begin{pmatrix} \alpha, & E_1, & \dots, & E_{n-1} \\ E_1, & c_{11}, & \dots, & c_{1,n-1} \\ \dots & \dots & \dots & \dots \\ E_{n-1}, & c_{n-1,1}, & \dots, & c_{n-1,n-1} \end{pmatrix} \pmod{2^e}$$

transformieren, in welcher α und im Falle $\omega_1 = 0$ auch eine der Größen $\alpha c_{ii} - E_i^2$ ungerade ist. Diese Form $\widehat{f}_{(1)}$ geht durch eine Substitution

$$S^{(1)} = \begin{pmatrix} 1, & K_1, & K_2, & \dots, & K_{n-1} \\ & 1, & 0, & \dots, & 0 \\ & & \dots & \dots & \dots \\ & & & & 1 \end{pmatrix},$$

in welcher die K_h den Kongruenzen

$$E_h + \alpha K_h \equiv 0 \pmod{2^e}$$

genügen, in eine Form

$$f_{(1)} \equiv \begin{pmatrix} \alpha, & 0, & \dots, & 0 \\ 0, & 2^{\omega_1} a_{11}^{(1)}, & \dots, & 2^{\omega_1} a_{1, n-1}^{(1)} \\ \dots & \dots & \dots & \dots \\ 0, & 2^{\omega_1} a_{n-1, 1}^{(1)}, & \dots, & 2^{\omega_1} a_{n-1, n-1}^{(1)} \end{pmatrix} \pmod{2^f}$$

über, in welcher $f^{(1)} = \{a_{ik}^{(1)}\}$ in bezug auf 2 primitiv ist. Hierin fällt, wenn $\omega_1 = 0$ ist, eine der Größen $\alpha 2^{\omega_1} a_{ii}^{(1)} \equiv a_{ii}^{(1)} \pmod{2}$ ungerade aus, und folglich wird in diesem Falle die Form $f^{(1)}$ eine erste Invariante σ gleich 1 besitzen.

2. An die Form $f_{(1)}$ können wir sofort dieselben Schlüsse anknüpfen wie in Absatz I, 2. dieses Kapitels, und wir gelangen, indem wir der Form $f^{(1)}$ $n - 2$ Zahlen $\omega_h^{(1)}$ zuerteilen und den Modul 2^f größer als die größte der Zahlen $2^{\hat{\omega}_1}, 2^{\hat{\omega}_2}, \dots, 2^{\hat{\omega}_{n-1}}$ wählen, in derselben Weise zum gewünschten Ziele:

$$\omega_h^{(1)} = \omega_h \quad (h > 1), \quad \partial_h = h\omega_1 + (h-1)\omega_2 + \dots + \omega_h, \quad \omega_h \geq 0.$$

$$\sigma_1 = 2. \quad \text{— 1. Jede in bezug auf 2 primitive Form } f = \sum_{i,k=1}^n a_{ik} x_i x_k,$$

welche eine Invariante $\sigma_1 = 2$ besitzt, ist einer Form

$$f_{(2)} \equiv \begin{pmatrix} 2\alpha, \mathfrak{A}, & 0, & \dots, & 0 \\ \mathfrak{A}, 2\tilde{\alpha}, & 0, & \dots, & 0 \\ 0, & 0, & 2^{\omega_2} a_{11}^{(2)}, & \dots, & 2^{\omega_2} a_{1, n-2}^{(2)} \\ \dots & \dots & \dots & \dots & \dots \\ 0, & 0, & 2^{\omega_2} a_{n-2, 1}^{(2)}, & \dots, & 2^{\omega_2} a_{n-2, n-2}^{(2)} \end{pmatrix} \pmod{2^f},$$

$$f_{(2)} \equiv 2(\alpha \xi^2 + \mathfrak{A} \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2) + 2^{\omega_2} \sum_{i,k=1}^{n-2} a_{ik}^{(2)} x_i^{(2)} x_k^{(2)} \pmod{2^f}$$

äquivalent, in welcher die Zahl $\mathfrak{A}$ einen willkürlich gewählten ungeraden Wert hat*), α ungerade und $f^{(2)} = \sum_{i,k=1}^{n-2} a_{ik}^{(2)} x_i^{(2)} x_k^{(2)}$ eine in bezug auf 2 primitive Form ist.

Beweis: Wegen $\sigma_1 = 2$ müssen alle a_{ii} gerade ausfallen, während mindestens eines der a_{ik} ($i \neq k$) ungerade wird. Wir können annehmen, daß gleichzeitig $a_{ik} \equiv 1 \pmod{2}$ und $a_{ii} \equiv 2 \pmod{4}$ gilt. Sind nämlich für ein ungerades a_{ik} gleichzeitig a_{ii} und a_{kk} durch 4 teilbar, so genügt es, auf f eine Substitution $S_{(i,k)}^{\pm 1}$ anzuwenden, um zu erzielen, daß der Koeffizient a_{ii} durch die Zahl $a_{ii} \pm 2a_{ik} + a_{kk} \equiv 2 \pmod{4}$ ersetzt wird.

*) Wir bezeichnen mit $\mathfrak{A}$ stets ungerade Zahlen, welche willkürlich gewählt werden dürfen.

Durch eine geeignete Umstellung der Variablen mittels der Substitutionen $P_{(1,i)}$, $P_{(2,k)}$ können wir daher der Form f die Gestalt verleihen

$$\hat{f}_{(2)} \equiv \begin{pmatrix} 2\alpha, & A, & E_1, & \dots, & E_{n-2} \\ A, & 2a, & \tilde{E}_1, & \dots, & \tilde{E}_{n-2} \\ E_1, & \tilde{E}_1, & c_{11}, & \dots, & c_{1,n-2} \\ \cdot & \cdot & \cdot & \cdot & \cdot \\ E_{n-2}, & \tilde{E}_{n-2}, & c_{n-2,1}, & \dots, & c_{n-2,n-2} \end{pmatrix} \pmod{2^f},$$

in welcher α und A ungerade sind. Auf dieses $\hat{f}_{(2)}$ wenden wir eine Substitution

$$S^{(2)} = \begin{pmatrix} 1, & M, & K_1, & K_2, & \dots, & K_{n-2} \\ & 1, & \tilde{K}_1, & \tilde{K}_2, & \dots, & \tilde{K}_{n-2} \\ & & 1, & 0, & \dots, & 0 \\ & & & \cdot & \cdot & \cdot \\ & & & & & 1 \end{pmatrix}$$

an. Eine solche Substitution läßt 2α ungeändert, während dieselbe

$$A \text{ in } A + 2\alpha M$$

und

$$E_h \text{ in } E_h + 2\alpha K_h + A\tilde{K}_h,$$

$$\tilde{E}_h \text{ in } (\tilde{E}_h + AK_h + 2a\tilde{K}_h) + M(E_h + 2\alpha K_h + A\tilde{K}_h)$$

verwandelt. Bezeichnet $\mathfrak{A}$ eine beliebige ungerade GröÙe, so hat die Kongruenz

$$2\alpha M \equiv \mathfrak{A} - A \pmod{2^f}$$

stets Lösungen, und ebenso sind wegen

$$4\alpha a - A^2 \equiv 1 \pmod{2}$$

die Kongruenzen

$$\begin{aligned} E_h + 2\alpha K_h + A\tilde{K}_h &\equiv 0 \\ \tilde{E}_h + AK_h + 2a\tilde{K}_h &\equiv 0 \end{aligned} \pmod{2^f}$$

auflösbar; sie ergeben

$$K_h \equiv \frac{A\tilde{E}_h - 2aE_h}{4\alpha a - A^2}, \quad \tilde{K}_h \equiv \frac{AE_h - 2\alpha\tilde{E}_h}{4\alpha a - A^2} \pmod{2^f}.$$

Mit Hilfe der so bestimmten Werte der M , K_h , $\tilde{K}_h$ geht daher die Form $\hat{f}_{(2)}$ in eine Form $f_{(2)}$ von der gewünschten Art über. In dieser Form $f_{(2)}$ müssen wegen $\sigma_1 = 2$ die sämtlichen $2^{\omega_2} a_{ii}^{(2)}$ gerade sein. Folglich sind im Falle $\omega_2 = 0$ auch alle $a_{ii}^{(2)}$ gerade, und die Form $f^{(2)} = \{a_{ik}^{(2)}\}$ wird in diesem Falle eine erste Invariante σ gleich 2 besitzen.

2. Die höchste Potenz von 2, welche allen $A_{(2)h}$ der Form $f_{(2)}$ ($\sim f$) gemeinsam ist, d. h. die höchste in d_{h-1} aufgehende Potenz von 2, sei wieder $2^{\partial_{h-1}}$, die höchste in den $A_h^{(2)}$ der Form $f^{(2)}$ aufgehende Potenz

von 2 sei $2^{\partial_h^{(2)}-1}$. Wir bemerken, daß die Potenz 2^{2^1} , welche unter anderem in $4\alpha\tilde{\alpha} - \mathfrak{U}^2$ aufgehen muß, offenbar gleich 1 ist und daß infolgedessen $\partial_1 = 0$ wird. Wir haben für die Größen $A_{(2)^h}$ die Kongruenzen

$$\begin{aligned} & 2\alpha \cdot 2^{(h-1)\omega_2} A_{h-1}^{(2)}, \\ A_{(2)^h} & \equiv (4\alpha\tilde{\alpha} - \mathfrak{U}^2) 2^{(h-2)\omega_2} A_{h-2}^{(2)}, \quad \mathfrak{U} \cdot 2^{(h-1)\omega_2} A_{h-1}^{(2)}, \quad 2^h \omega_2 A_h^{(2)}, \quad 0 \pmod{2^t}. \\ & 2\tilde{\alpha} \cdot 2^{(h-1)\omega_2} A_{h-1}^{(2)}, \end{aligned}$$

Wählen wir nun t größer als die größte der Zahlen $\partial_1, \partial_2, \dots, \partial_{n-1}$, so wird die höchste in allen $A_{(2)^h}$ aufgehende Potenz von 2, nämlich 2^{∂_h-1} , zugleich die höchste Potenz von 2 sein, welche allen auf der rechten Seite der vorstehenden Kongruenzen befindlichen Größen gemeinsam ist. Es ist aber die höchste in allen

$$(4\alpha\tilde{\alpha} - \mathfrak{U}^2) 2^{(h-2)\omega_2} A_{h-2}^{(2)},$$

in allen

$$2\alpha \cdot 2^{(h-1)\omega_2} A_{h-1}^{(2)}, \quad \mathfrak{U} \cdot 2^{(h-1)\omega_2} A_{h-1}^{(2)}, \quad 2\tilde{\alpha} \cdot 2^{(h-1)\omega_2} A_{h-1}^{(2)},$$

in allen

$$2^h \omega_2 A_h^{(2)}$$

aufgehende Potenz von 2 resp.

$$2^{(h-2)\omega_2 + \partial_{h-3}^{(2)}}, \quad 2^{(h-1)\omega_2 + \partial_{h-2}^{(2)}}, \quad 2^h \omega_2 + \partial_{h-1}^{(2)}.$$

Sonach wird ∂_{h-1} mit der kleinsten der drei Zahlen

$$(h-2)\omega_2 + \partial_{h-3}^{(2)}, \quad (h-1)\omega_2 + \partial_{h-2}^{(2)}, \quad h\omega_2 + \partial_{h-1}^{(2)}$$

übereinstimmen müssen.

Nehmen wir jetzt an, für die Form $f^{(2)} = \{a_{ik}^{(2)}\}$ von $n-2$ Variablen wäre der Satz C. richtig, so werden die Größen $\omega_h^{(2)}$, welche durch die Gleichungen

$$\partial_h^{(2)} = h\omega_1^{(2)} + (h-1)\omega_2^{(2)} + \dots + \omega_h^{(2)}$$

bestimmt sind, nicht negativ ausfallen, und es müssen demnach die Ungleichungen

$$\partial_{h-3}^{(2)} \leq \partial_{h-2}^{(2)} \leq \partial_{h-1}^{(2)}$$

statthaben. Umsomehr gelten dann wegen $\omega_2 \geq 0$ die weiteren Ungleichungen

$$(h-2)\omega_2 + \partial_{h-3}^{(2)} \leq (h-1)\omega_2 + \partial_{h-2}^{(2)} \leq h\omega_2 + \partial_{h-1}^{(2)}.$$

Mithin wird ∂_{h-1} , welches der kleinsten der vorstehenden drei Zahlen gleich werden soll, gleich $(h-2)\omega_2 + \partial_{h-3}^{(2)}$ werden, und wir erhalten in diesem letzten Fall, indem wir

$$\omega_1 = \partial_1 = 0 \quad \text{und} \quad \omega_{h-2}^{(2)} = \omega_h \quad (h > 2)$$

setzen, gleichfalls

$$\partial_h = h\omega_1 + (h-1)\omega_2 + \dots + \omega_h, \quad \omega_h \geq 0.$$

Wir führen jetzt einige Bezeichnungen ein, welche für die folgenden Untersuchungen wichtig sind. Wir setzen für eine jede Primzahl q ($q = p$ oder $= 2$)

$$v_1 = \omega_1, v_2 = \omega_1 + \omega_2, \dots, v_{n-1} = \omega_1 + \omega_2 + \dots + \omega_{n-1},$$

woraus sich sofort

$$\partial_1 = v_1, \partial_2 = v_1 + v_2, \dots, \partial_{n-1} = v_1 + v_2 + \dots + v_{n-1}$$

ergibt. Die Potenzen q^{ω_h} sind diejenigen Potenzen von q , welche in den Größen $\frac{d_h}{d_{h-1}}$ aufgehen. Ferner nehmen wir an, von den $n-1$ Zahlen ω_h seien im ganzen $\lambda-1$ ($1 \leq \lambda \leq n$) von Null verschieden, nämlich die Größen

$$(\vartheta_0 = 0) \quad \omega_{\vartheta_1}, \omega_{\vartheta_2}, \dots, \omega_{\vartheta_{\lambda-2}}, \omega_{\vartheta_{\lambda-1}}, \quad (\vartheta_{\lambda} = n)$$

und wir bestimmen λ Zahlen κ_k durch die Gleichungen

$$\vartheta_1 = \kappa_1, \vartheta_2 = \kappa_1 + \kappa_2, \dots, n = \vartheta_{\lambda} = \kappa_1 + \kappa_2 + \dots + \kappa_{\lambda}, (\kappa_k = \vartheta_k - \vartheta_{k-1}).$$

So oft der besondere Wert der Primzahl q hervorgehoben werden soll, werden wir den sämtlichen Größen $\omega, v, \partial, \lambda$ die Zahl q in Klammern beifügen.

Kap. III. Hauptformenreste und Hauptrepräsentanten für einen Modul N .

I. Der Modul N möge ein Produkt von c_0 zueinander relativ primen Zahlen $N_1, N_2, \dots, N_{c_0}$ sein. Wir beweisen den folgenden Satz:

Wenn in der Klasse f sich c_0 Formen $R_c = \bar{S}_c f S_c$ ($c = 1, 2, \dots, c_0$) vorfinden, welche die Reste

$$R_c \equiv \{r_{ik}^{(c)}\} \pmod{N_c} \quad (c = 1, 2, \dots, c_0)$$

besitzen, so können wir einen Repräsentanten

$$R = \bar{S} f S \equiv \{r_{ik}\} \pmod{N}$$

dieser Klasse angeben, welcher den Kongruenzen

$$r_{ik} \equiv r_{ik}^{(c)} \pmod{N_c} \quad (c = 1, 2, \dots, c_0)$$

genügt.

Beweis: S_c läßt sich als Substitution von der Determinante 1 bekanntlich in eine gewisse Anzahl Teilsubstitutionen von der Art

$$x_i = x'_i + M_c x'_k, \quad x_k = x'_k, \quad x_h = x'_h \quad (h \neq i, k)$$

zerlegen. Führen wir in diesen Teilsubstitutionen anstatt der M_c beliebige andere Zahlen M_c^* ein, welche gleichzeitig den Kongruenzen

$$M_c^* \equiv M_c \pmod{N_c} \quad \text{und} \quad M_c^* \equiv 0 \pmod{\frac{N}{N_c}}$$

genügen, und setzen die so veränderten Teilsubstitutionen in genau derselben Weise, wie sie durch Zerlegung von S_c entstanden sind, zu einer neuen Substitution S_c^* zusammen, so wird offenbar

$$S_c^* \equiv S_c \pmod{N_c}, \quad S_c^* \equiv \begin{pmatrix} 1, 0, \dots, 0 \\ 0, 1, \dots, 0 \\ \dots \dots \dots \\ 0, 0, \dots, 1 \end{pmatrix} \pmod{\frac{N}{N_c}}.$$

Es folgt also

$$R_c^* = \overline{S_c^*} f S_c^* \equiv R_c \pmod{N_c}, \quad R_c^* \equiv f \pmod{\frac{N}{N_c}},$$

und für die Form

$$R = \overline{S_1^* S_2^* \dots S_{c_0}^*} \cdot f \cdot S_1^* S_2^* \dots S_{c_0}^*$$

haben wir in der Tat

$$R \equiv \overline{S_c^*} f S_c^* \equiv R_c \pmod{N_c}.$$

Dieser Satz zeigt uns die Möglichkeit, vermittels der Formenreste nach den Teilmoduln $N_1, N_2, \dots, N_{c_0}$ einen Formenrest nach dem zusammengesetzten Modul N zu finden. Insbesondere können die N_c die verschiedenen in N enthaltenen Primzahlpotenzen bedeuten; wir kommen daher mit der Bestimmung ausgezeichneter Formenreste für Primzahlpotenzen q^t als Moduln aus. Es genügt auch völlig, wenn wir nur Moduln q^t oberhalb gewisser Grenzen betrachten; denn es ist ein jeder Formenrest für einen Modul N zugleich Formenrest für alle in N aufgehenden Moduln N_0 ; insbesondere wird daher ein jeder Formenrest für einen Modul q^t auch Formenrest für alle Moduln q^s sein, welche kleiner als q^t sind.

Durch eine wiederholte Anwendung der in Kap. II gewonnenen Sätze erkennen wir, daß eine jede Formenklasse für jeden Modul q^t Reste besitzt, welche sich als Summen von Formen mit einer oder mit zwei Variablen darstellen.

II. ($q = p$.) *Aus der Klasse äquivalenter Formen, welcher die zu p primitive Form f angehört, kann ein Repräsentant*

$$\varphi \equiv \begin{pmatrix} \alpha_1, \\ p^{\omega_1} \alpha_2, \\ p^{\omega_1 + \omega_2} \alpha_3, \\ \dots, \\ p^{\omega_1 + \omega_2 + \dots + \omega_{n-1}} \alpha_n \end{pmatrix} \pmod{p^t}$$

ausgewählt werden, in welchem die $\alpha_1, \alpha_2, \alpha_3, \dots, \alpha_n$ sämtlich zu p relativ prime Zahlen sind.

Wir beweisen diesen Satz, welcher für $n = 1$ selbstverständlich ist, vermittels eines Schlusses von $n - 1$ auf n . Angenommen dieser Satz sei für Formen mit $n - 1$ Variablen bereits bewiesen, so können wir die Form $p^{\omega_1} f^{(1)}$ des Repräsentanten

$$f_{(1)} \equiv \alpha \xi^2 + p^{\omega_1} f^{(1)} \pmod{p^t}$$

durch eine gewisse Substitution in eine andere überführen, deren Rest nach dem Modul p^t aus lauter eingliedrigen Einzelformen besteht. Durch

dieselbe Substitution geht dann $f_{(1)}$ in eine Form von der Gestalt φ über, womit unser Satz bewiesen ist.

Mit Benutzung der am Ende von Kap. II eingeführten Bezeichnungen können wir für φ schreiben

$$\varphi \equiv \varphi_{(1)} + \varphi_{(2)} + \cdots + \varphi_{(\lambda)} \equiv \sum_{k=1}^{\lambda} \varphi_{(k)} \pmod{p^t},$$

$$\varphi_{(k)} \equiv p^{v_{\mathfrak{P}_{k-1}}} \sum_{s=1}^{x_k} \alpha_s^{(k)} \xi_s^{(k)} \xi_s^{(k)} \pmod{p^t}, \quad (\alpha_s^{(k)} = \alpha_{\mathfrak{P}_{k-1}+s})$$

$$0 < v_{\mathfrak{P}_1} < v_{\mathfrak{P}_2} < \cdots < v_{\mathfrak{P}_{\lambda-1}}.$$

($q = 2$.) 1. Es sei f eine in bezug auf 2 primitive Form, und es mögen die $x_1 - 1$ ersten Zahlen ω_h dieser Form gleich Null sein ($\omega_1 = 0$, $\omega_2 = 0, \dots, \omega_{x_1-1} = 0$), während die x_1 te Zahl ω_{x_1} ($1 \leq x_1 \leq n$) die erste der Zahlen ω_h sei, welche nicht verschwindet. Wenn dann die Invariante σ_1 von f gleich 2 ist, muß $x_1 \equiv 0 \pmod{2}$ sein. Wir können, falls $\sigma_1 = 1$ ist, statt f einen Repräsentanten

$$f_{(x_1)} \equiv \begin{pmatrix} \alpha_1^{(1)}, & & & & & \\ & \alpha_2^{(1)}, & & & & \\ & & \dots, & & & \\ & & & \alpha_{x_1}^{(1)}, & & \\ & & & & 2^{\omega_{x_1}} a_{11}^{(x_1)}, \dots, 2^{\omega_{x_1}} a_{1, n-x_1}^{(x_1)} & \\ & & & & \dots & \\ & & & & 2^{\omega_{x_1}} a_{n-x_1, 1}^{(x_1)}, \dots, 2^{\omega_{x_1}} a_{n-x_1, n-x_1}^{(x_1)} \end{pmatrix} \pmod{2^t}$$

eingeführen, in welchem die Zahlen $\alpha_1^{(1)}, \alpha_2^{(1)}, \dots, \alpha_{x_1}^{(1)}$ ungerade sind und

$$f_{(x_1)} = \sum_{i, k=1}^{n-x_1} a_{ik}^{(x_1)} x_i^{(x_1)} x_k^{(x_1)}$$

eine in bezug auf 2 primitive Form vorstellt; falls aber $\sigma_1 = 2$ ist, einen Repräsentanten

$$f_{(x_1)} \equiv \begin{pmatrix} 2\alpha_1^{(1)}, \mathfrak{A}_1^{(1)} \\ \mathfrak{A}_1^{(1)}, 2\tilde{\alpha}_1^{(1)}, \\ & 2\alpha_2^{(1)}, \mathfrak{A}_2^{(1)} \\ & \mathfrak{A}_2^{(1)}, 2\tilde{\alpha}_2^{(1)}, \\ & & \dots, \\ & & 2\alpha_{\frac{x_1}{2}}^{(1)}, \mathfrak{A}_{\frac{x_1}{2}}^{(1)} \\ & & \mathfrak{A}_{\frac{x_1}{2}}^{(1)}, 2\tilde{\alpha}_{\frac{x_1}{2}}^{(1)}, \\ & & & \dots, \\ & & & 2^{\omega_{x_1}} a_{11}^{(x_1)}, \dots, 2^{\omega_{x_1}} a_{1, n-x_1}^{(x_1)} \\ & & & \dots & \\ & & & 2^{\omega_{x_1}} a_{n-x_1, 1}^{(x_1)}, \dots, 2^{\omega_{x_1}} a_{n-x_1, n-x_1}^{(x_1)} \end{pmatrix} \pmod{2^t},$$

in welchem die Zahlen $\mathfrak{A}_1^{(1)}, \mathfrak{A}_2^{(1)}, \dots, \mathfrak{A}_{\frac{\kappa_1}{2}}^{(1)}$ beliebige ungerade Werte haben, die Zahlen $\alpha_1^{(1)}, \alpha_2^{(1)}, \dots, \alpha_{\frac{\kappa_1}{2}}^{(1)}$ ungerade sind und die Form

$$f^{(\kappa_1)} = \sum_{i, k=1}^{n-\kappa_1} \alpha_{ik}^{(\kappa_1)} x_i^{(\kappa_1)} x_k^{(\kappa_1)} \text{ in bezug auf 2 primitiv ist.}$$

Beweis. Wir hatten in Kap. II für f im Falle $\sigma_1 = 1$ zunächst einen Repräsentanten

$$f_{(1)} \equiv \alpha \xi^2 + 2\omega_1 f^{(1)} \pmod{2^\epsilon} \quad [\alpha \equiv 1 \pmod{2}]$$

gewonnen, in welchem $f^{(1)} = \{a_{ik}^{(1)}\}$ eine in bezug auf 2 primitive Form vorstellte, die, falls $\omega_1 = 0$ war, eine erste Invariante σ gleich 1 besaß, und im Falle $\sigma_1 = 2$ einen Repräsentanten

$$f_{(2)} \equiv 2(\alpha \xi^2 + \mathfrak{A} \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2) + 2\omega_2 f^{(2)} \pmod{2^\epsilon}, \\ [\alpha \equiv 1, \mathfrak{A} \equiv 1 \pmod{2}],$$

in welchem $f^{(2)} = \{a_{ik}^{(2)}\}$ eine in bezug auf 2 primitive Form vorstellte, die, falls $\omega_2 = 0$ war, eine erste Invariante σ gleich 2 besaß.

Im Falle $\kappa_1 \leq \sigma_1$ haben diese Repräsentanten die Gestalt, welche unser Satz verlangt. Wenn $\kappa_1 > \sigma_1$ wird und mithin $\omega_{\sigma_1} = 0$ ist, besitzt die Form $f^{(\sigma_1)} = \{a_{ik}^{(\sigma_1)}\}$ ebenso wie f eine erste Invariante σ gleich σ_1 . Nun ergibt sich aus den in Kap. II bewiesenen Gleichungen $\omega_{h-\sigma_1}^{(\sigma_1)} = \omega_h$ ($h > \sigma_1$) infolge unserer Annahme über die Größen ω_h sofort

$$\omega_1^{(\sigma_1)} = 0, \omega_2^{(\sigma_1)} = 0, \dots, \omega_{\kappa_1 - \sigma_1 - 1}^{(\sigma_1)} = 0, \omega_{\kappa_1 - \sigma_1}^{(\sigma_1)} = \omega_{\kappa_1} > 0.$$

Folglich wird für die Form $f^{(\sigma_1)}$ schon die $(\kappa_1 - \sigma_1)^{\text{te}}$ der Zahlen $\omega_h^{(\sigma_1)}$ größer als Null ausfallen. Machen wir also die Annahme, unser Satz wäre bereits bewiesen, sobald die erste von 0 verschiedene Zahl ω sich unter den $\kappa_1 - \sigma_1$ ersten dieser Zahlen befindet, so werden wir $2^{\omega_{\sigma_1}} f^{(\sigma_1)}$ durch gewisse Substitutionen in eine Form überführen können, welche die in dem Satze geforderte Gestalt besitzt, und wenn $\sigma_1(f^{(\sigma_1)}) = \sigma_1(f) = 2$ ist, wird $\kappa_1 - \sigma_1 = \kappa_1 - 2 \equiv 0 \pmod{2}$ sein. Dieselben Substitutionen transformieren alsdann die Form $f_{(\sigma_1)}$ in Formen von der Gestalt $f_{(\kappa_1)}$, und im Falle $\sigma_1 = 2$ wird $\kappa_1 \equiv 0 \pmod{2}$ sein. Hierdurch ist unser Satz für den Fall bewiesen, daß eine der κ_1 ersten Zahlen ω nicht verschwindet.

In derselben Weise, in welcher der Form f die $n - 1$ Zahlen ω_h entsprechen, gehören zur Form $f^{(\kappa_1)}$ gewisse $n - \kappa_1 - 1$ Größen $\omega_h^{(\kappa_1)}$; durch eine wiederholte Anwendung der Relationen

$$\omega_{h-\sigma_1}^{(\sigma_1)} = \omega_h \quad (h > \sigma_1)$$

finden wir leicht

$$\omega_h = \omega_{h-\kappa_1}^{(\kappa_1)} \quad (h > \kappa_1).$$

2. Eine in bezug auf 2 primitive Form f kann in einen Repräsentanten

$$\varphi \equiv \varphi_{(1)} + \varphi_{(2)} + \cdots + \varphi_{(\lambda)} \equiv \sum_{k=1}^{\lambda} \varphi_{(k)} \pmod{2^t}$$

transformiert werden, in welchem die Formen $\varphi_{(k)}$ Reste von der Gestalt

$$(R_I) \quad \varphi_{(k)} \equiv 2^{v_{\varphi_{k-1}}} \sum_{s=1}^{x_k} \alpha_s^{(k)} \xi_s^{(k)} \xi_s^{(k)} \pmod{2^t}$$

oder im Falle $x_k \equiv 0 \pmod{2}$ auch von der Gestalt

$$(R_{II}) \quad \varphi_{(k)} \equiv 2^{v_{\varphi_{k-1}} + 1} \sum_{s=1}^{\frac{1}{2} x_k} (\alpha_s^{(k)} \xi_s^{(k)} \xi_s^{(k)} + \mathfrak{A}_s^{(k)} \xi_s^{(k)} \tilde{\xi}_s^{(k)} + \tilde{\alpha}_s^{(k)} \tilde{\xi}_s^{(k)} \xi_s^{(k)}) \pmod{2^t}$$

besitzen.

$$0 < v_{\varphi_1} < v_{\varphi_2} < \cdots < v_{\varphi_{\lambda-1}}.$$

Für ein $\lambda = 1$ erhellt die Richtigkeit dieses Satzes schon aus den in 1. gewonnenen Resultaten. Infolge der Relationen $\omega_h = \omega_{h-x_1}^{(x_1)}$ ($h > x_1$) erkennen wir, daß, wenn $\lambda - 1$ von den $n - 1$ zur Form $f_{(x_1)}$ gehörigen Zahlen ω_h nicht verschwinden, alsdann von den $n - x_1 - 1$ zur Form $f^{(x_1)}$ gehörigen Zahlen $\omega_h^{(x_1)}$ nur $\lambda - 2$ größer als Null ausfallen. Machen wir daher die Annahme, unser Satz wäre richtig für alle Formen, welche nur $\lambda - 2$ von 0 verschiedene Zahlen ω besitzen, so wird die Form $2^{\omega_{x_1}} f^{(x_1)}$ für den Modul 2^t durch eine gewisse Substitution auf eine Gestalt gebracht werden können, wie sie unser Satz verlangt. Durch die nämliche Substitution geht dann die Form $f_{(x_1)}$ wirklich in eine Form von der Gestalt φ über.

III. Jeder Formenrest einer Klasse f , welcher die Gestalt $\varphi \pmod{q^t}$ hat, soll ein *Hauptformenrest* [[oder Hauptrest]] dieser Klasse für den Modul q^t heißen, während jede Form der Klasse f , welche einen Hauptrest für den Modul q^t liefert, ein *Hauptrepräsentant* dieser Klasse für den Modul q^t genannt werden möge.

In einem Hauptrest für einen Modul q^t sind entweder die seitlichen Koeffizienten alle $\equiv 0 \pmod{q^t}$ oder können im Falle $q = 2$ zum Teil (mit einer leichten Einschränkung) nach dem Modul q^t willkürlich gewählt werden. Jedenfalls kann man es in zwei verschiedenen Hauptresten für einen Modul q^t stets ohne Mühe erreichen, daß diese Koeffizienten modulo q^t dieselben Werte annehmen, und auf diese Weise wird ein Hauptrest für einen Modul q^t nur von den Resten seiner n mittleren Koeffizienten abhängen.

Eine Form φ soll ein Hauptrepräsentant für einen zusammengesetzten Modul N und ihr Rest soll ein Hauptrest für den Modul N heißen, sobald φ einen Hauptrepräsentanten für alle in N aufgehenden Primzahlpotenzen vorstellt.

Kap. IV. Bedingungen für die Existenz einer Ordnung $\mathcal{O}$.

Mit Hilfe der Hauptformenreste für Modul 2^t können wir jetzt Bedingungen finden, an welche die Existenz einer Ordnung

$$\begin{pmatrix} \sigma_1, \sigma_2, \dots, \sigma_{n-1} \\ o_1, o_2, \dots, o_{n-1} \end{pmatrix}, \quad I$$

gebunden ist.

I. Es sei f eine in bezug auf 2 primitive Form.

1. Wir gehen zunächst von den in Kap. III aufgestellten Repräsentanten

$$f_{(\kappa_1)} \equiv \Phi_{(1)} + 2^{\omega_{\kappa_1}} f^{(\kappa_1)} \pmod{2^t}$$

aus, in denen $\Phi_{(1)}$ und $f^{(\kappa_1)}$ in bezug auf 2 primitive Formen bezeichnen. Die $\kappa_1 - 1$ Zahlen σ der Form $\Phi_{(1)}$ setzen wir gleich $\varrho_1^{(1)}, \varrho_2^{(1)}, \dots, \varrho_{\kappa_1-1}^{(1)}$ und die $n - \kappa_1 - 1$ Zahlen σ der Form $f^{(\kappa_1)}$ gleich $\sigma_1^{(\kappa_1)}, \sigma_2^{(\kappa_1)}, \dots, \sigma_{n-\kappa_1-1}^{(\kappa_1)}$. Ferner bezeichnen wir die h -reihigen Unterdeterminanten der Formen $f_{(\kappa_1)}, f^{(\kappa_1)}, \Phi_{(1)}$ durch $F_{(h; \kappa_1)}$ oder $P_{(h; \kappa_1)}$, $F_h^{(\kappa_1)}$ oder $P_h^{(\kappa_1)}$, $D_h^{(1)}$ oder $M_h^{(1)}$, indem wir einen Buchstaben F, D anwenden, sobald die betreffende Unterdeterminante symmetrisch ist, dagegen einen Buchstaben P, M , sobald dieselbe unsymmetrisch ausfällt. Die aus den ersten κ_1 Reihen von $f_{(\kappa_1)}$ gebildete symmetrische Determinante möge gleich $|\Phi_{(1)}|$ sein.

Für die Invarianten σ der Form $f_{(\kappa_1)}$ ($\sim f$) erhalten wir die folgenden Sätze.

Es ist

$$\sigma_h = \varrho_h^{(1)} \quad (h < \kappa_1).$$

In der Tat zeigen die Kongruenzen

$$\begin{aligned} F_{(h; \kappa_1)} &\equiv D_h^{(1)}, \quad 0 \\ P_{(h; \kappa_1)} &\equiv M_h^{(1)}, \quad 0 \pmod{2} \quad [h < \kappa_1], \end{aligned}$$

daß einerseits die höchsten Potenzen von 2, welche in allen Zahlen $F_{(h; \kappa_1)}, P_{(h; \kappa_1)}$ und in allen $D_h^{(1)}, M_h^{(1)}$ aufgehen, und andererseits die größten Potenzen, welche in allen $F_{(h; \kappa_1)}, 2P_{(h; \kappa_1)}$ und in allen $D_h^{(1)}, 2M_h^{(1)}$ aufgehen, nach dem Modul 2 kongruent sind. Da nun die höchste in den Größen $F_{(h; \kappa_1)}, P_{(h; \kappa_1)}$ ($h < \kappa_1$) aufgehende Potenz von 2, nämlich $2^{\partial_{h-1}}$, infolge unserer Annahme über die Zahlen ω_h ($h < \kappa_1$) gleich 1 ist, so wird die höchste in allen $D_h^{(1)}, M_h^{(1)}$ aufgehende Potenz von 2 ebenfalls gleich 1 sein müssen, und die höchsten Potenzen von 2, welche in den $F_{(h; \kappa_1)}, 2P_{(h; \kappa_1)}$ und in den $D_h^{(1)}, 2M_h^{(1)}$ aufgehen, werden die Werte σ_h und $\varrho_h^{(1)}$ haben. Jetzt müssen die Größen σ_h und $\varrho_h^{(1)}$, da sie nach dem Modul 2 kongruent und zugleich ≤ 2 sind, identisch sein, und es ergibt sich demnach in der Tat $\sigma_h = \varrho_h^{(1)}$ für $h < \kappa_1$.

Es ist

$$\sigma_h = 1 \quad (h = \kappa_1).$$

Die Größe σ_{κ_1} muß in der Zahl $|\Phi_{(1)}|$ aufgehen; diese ist aber offenbar ungerade, mithin ist $\sigma_{\kappa_1} = 1$.

Es ist

$$\sigma_h = \sigma_{h-\kappa_1}^{(\kappa_1)} \quad (h > \kappa_1).$$

Für die Determinanten $F_{(h; \kappa_1)}$ und $P_{(h; \kappa_1)}$ ($h > \kappa_1$) bestehen die Kongruenzen

$$F_{(h; \kappa_1)} \equiv |\Phi_{(1)}| \cdot 2^{(h-\kappa_1)\omega_{\kappa_1}} F_{h-\kappa_1}^{(\kappa_1)}, \quad 2^{(h-h_0)\omega_{\kappa_1}} D_{h-h_0}^{(1)} F_{h-h_0}^{(\kappa_1)}, \quad 0 \pmod{2^t},$$

$$P_{(h; \kappa_1)} \equiv |\Phi_{(1)}| \cdot 2^{(h-\kappa_1)\omega_{\kappa_1}} P_{h-\kappa_1}^{(\kappa_1)}, \quad \begin{cases} 2^{(h-h_0)\omega_{\kappa_1}} D_{h-h_0}^{(1)} P_{h-h_0}^{(\kappa_1)}, \\ 2^{(h-h_0)\omega_{\kappa_1}} M_{h-h_0}^{(1)} F_{h-h_0}^{(\kappa_1)}, \\ 2^{(h-h_0)\omega_{\kappa_1}} M_{h-h_0}^{(1)} P_{h-h_0}^{(\kappa_1)}, \end{cases} \quad 0 \pmod{2^t}$$

$$(h_0 = 0, 1, \dots, \kappa_1 - 1; D_0^{(1)} = 1, M_0^{(1)} = 0).$$

Machen wir jetzt die Annahme $t > \partial_{h-1} + 1$, so wird eine jede Größenreihe, welche nach dem Modul 2^t mit der Größenreihe $F_{(h; \kappa_1)}$, $2P_{(h; \kappa_1)}$ übereinstimmt, durch eben dieselbe Potenz $\sigma_h 2^{\partial_h-1}$ und durch keine höhere Potenz von 2 teilbar sein wie die Größenreihe $F_{(h; \kappa_1)}$, $2P_{(h; \kappa_1)}$. Nun ist die höchste Potenz von 2, welche den Zahlen

$$|\Phi_{(1)}| \cdot 2^{(h-\kappa_1)\omega_{\kappa_1}} F_{h-\kappa_1}^{(\kappa_1)}, \quad 2|\Phi_{(1)}| \cdot 2^{(h-\kappa_1)\omega_{\kappa_1}} P_{h-\kappa_1}^{(\kappa_1)}$$

gemeinsam ist, gleich

$$\sigma_{h-\kappa_1}^{(\kappa_1)} \cdot 2^{(h-\kappa_1)\omega_{\kappa_1} + \partial_{h-\kappa_1-1}^{(\kappa_1)}},$$

während in den Zahlen

$$\begin{aligned} & 2 \cdot 2^{(h-h_0)\omega_{\kappa_1}} D_{h-h_0}^{(1)} P_{h-h_0}^{(\kappa_1)}, \\ & 2^{(h-h_0)\omega_{\kappa_1}} D_{h-h_0}^{(1)} F_{h-h_0}^{(\kappa_1)}, \quad 2 \cdot 2^{(h-h_0)\omega_{\kappa_1}} M_{h-h_0}^{(1)} F_{h-h_0}^{(\kappa_1)}, \\ & 2 \cdot 2^{(h-h_0)\omega_{\kappa_1}} M_{h-h_0}^{(1)} P_{h-h_0}^{(\kappa_1)}, \end{aligned}$$

eine Potenz

$$\begin{aligned} & \geq 2^{(h-h_0)\omega_{\kappa_1} + \partial_{h-h_0-1}^{(\kappa_1)}} \geq 2^{\omega_{\kappa_1}} \cdot 2^{(h-\kappa_1)\omega_{\kappa_1} + \partial_{h-\kappa_1-1}^{(\kappa_1)}} \\ & \geq \sigma_{h-\kappa_1}^{(\kappa_1)} \cdot 2^{(h-\kappa_1)\omega_{\kappa_1} + \partial_{h-\kappa_1-1}^{(\kappa_1)}} \end{aligned}$$

aufgeht.

Wir bekommen daher für ein $h > \kappa_1$ die Beziehung

$$\sigma_{h-\kappa_1}^{(\kappa_1)} \cdot 2^{(h-\kappa_1)\omega_{\kappa_1} + \partial_{h-\kappa_1-1}^{(\kappa_1)}} = \sigma_h \cdot 2^{\partial_h-1}.$$

Aus der oben gewonnenen Gleichung $\omega_{h-\kappa_1}^{(\kappa_1)} = \omega_h$ ($h > \kappa_1$) folgt aber ohne Schwierigkeit

$$2^{(h-\kappa_1)\omega_{\kappa_1} + \partial_{h-\kappa_1-1}^{(\kappa_1)}} = 2^{\partial_h-1},$$

und wir gelangen sonach wirklich zu dem Resultat $\sigma_h = \sigma_{h-\kappa_1}^{(\kappa_1)}$ ($h > \kappa_1$).

2. Sei jetzt

$$\varphi \equiv \Phi_{(1)} + 2^{\omega_{\mathfrak{I}_1}} \{ \Phi_{(2)} + 2^{\omega_{\mathfrak{I}_2}} [\Phi_{(3)} + \dots + 2^{\omega_{\mathfrak{I}_{\lambda-1}}} (\Phi_{(\lambda)}) \cdot \cdot] \} \pmod{2^f}$$

ein Hauptrepräsentant der Klasse f nach dem Modul 2^f . Wir wollen die Invarianten σ der Formen $\Phi_{(k)}$ durch

$$\varrho_1^{(k)}, \varrho_2^{(k)}, \dots, \varrho_{\kappa_k-2}^{(k)}, \varrho_{\kappa_k-1}^{(k)}$$

bezeichnen.

D. Die Invarianten σ_h können vermittels der Relationen

$$\sigma_{\mathfrak{I}_{k-1}+h} = \varrho_h^{(k)} \quad (h = 1, 2, \dots, \kappa_k - 1)$$

und

$$\sigma_{\mathfrak{I}_k} = 1$$

durch die Invarianten $\varrho_h^{(k)}$ ausgedrückt werden.

Beweis: In 1. erhielten wir die Gleichungen

$$\sigma_h = \varrho_h^{(1)} \quad (h < \kappa_1), \quad \sigma_{\mathfrak{I}_1} = 1, \quad \sigma_h = \sigma_{h-\kappa_1}^{(\kappa_1)} \quad (h > \kappa_1);$$

dieselben rechtfertigen im Falle $\lambda = 1$ den Satz D. vollständig, während sie uns in den Fällen $\lambda > 1$ eine Gelegenheit zur Anwendung eines Schlusses von $\lambda - 1$ auf λ verschaffen. Wie wir in Kap. III gesehen haben, gehört der Form $f^{(\kappa_1)}$ in derselben Weise die Zahl $\lambda - 1$ zu, wie der Form f die Zahl λ entspricht, und es ist sofort klar, daß der Hauptrepräsentant $\varphi \pmod{2^f}$ der Klasse f einen Hauptrepräsentanten

$$2^{\omega_{\kappa_1}} \varphi^{(\kappa_1)} \equiv 2^{\omega_{\kappa_1}} \{ \Phi_{(2)} + 2^{\omega_{\mathfrak{I}_2}} [\Phi_{(3)} + \dots + 2^{\omega_{\mathfrak{I}_{\lambda-1}}} (\Phi_{(\lambda)}) \cdot \cdot] \} \pmod{2^f}$$

für die Form $2^{\omega_{\kappa_1}} f^{(\kappa_1)}$ mit sich führt. Machen wir nun die Annahme, der Satz sei richtig für Formen mit $\lambda - 2$ von Null verschiedenen Zahlen ω , so erhalten wir

$$\sigma_{\mathfrak{I}_{k-1}-\kappa_1+h}^{(\kappa_1)} = \varrho_h^{(k)} \quad (k > 1; \quad h = 1, \dots, \kappa_k - 1)$$

und

$$\sigma_{\mathfrak{I}_k-\kappa_1}^{(\kappa_1)} = 1 \quad (k > 1),$$

woraus auf der Stelle die Gleichungen

$$\sigma_{\mathfrak{I}_{k-1}+h} = \varrho_h^{(k)} \quad (k > 1; \quad h = 1, \dots, \kappa_k - 1)$$

und

$$\sigma_{\mathfrak{I}_k} = 1 \quad (k > 1)$$

hervorgehen; dieselben beweisen zusammen mit den Gleichungen

$$\sigma_h = \varrho_h^{(1)} \quad (h = 1, \dots, \kappa_1 - 1) \quad \text{und} \quad \sigma_{\mathfrak{I}_1} = 1$$

den Satz D. für die Form f , welche $\lambda - 1$ von Null verschiedene Zahlen ω besitzt.

3. Je nachdem $\Phi_{(k)}$ einen Rest von der Gestalt

$$(R_I) \quad \Phi_{(k)} \equiv \begin{pmatrix} \alpha_1^{(k)}, & & & \\ & \alpha_2^{(k)} & & \\ & & \dots & \\ & & & \alpha_{\kappa_k}^{(k)} \end{pmatrix} \pmod{2^{t-v_{\vartheta_{k-1}}}}$$

oder von der Gestalt

$$(R_{II}) \quad \Phi_{(k)} \equiv \begin{pmatrix} 2\alpha_1^{(k)}, \mathfrak{A}_1^{(k)}, \\ \mathfrak{A}_1^{(k)}, 2\tilde{\alpha}_1^{(k)}, & & & \\ & 2\alpha_2^{(k)}, \mathfrak{A}_2^{(k)}, \\ & \mathfrak{A}_2^{(k)}, 2\tilde{\alpha}_2^{(k)}, & & \\ & & \dots & \\ & & & 2\alpha_{\frac{1}{2}\kappa_k}^{(k)}, \mathfrak{A}_{\frac{1}{2}\kappa_k}^{(k)}, \\ & & & \mathfrak{A}_{\frac{1}{2}\kappa_k}^{(k)}, 2\tilde{\alpha}_{\frac{1}{2}\kappa_k}^{(k)} \end{pmatrix} \pmod{2^{t-v_{\vartheta_{k-1}}}}$$

läßt, erhalten die Größen $\varrho_h^{(k)}$, wie wir uns leicht überzeugen, die Werte

$$(R_I) \quad \varrho_1^{(k)} = 1, \varrho_2^{(k)} = 1, \dots, \varrho_{\kappa_k-2}^{(k)} = 1, \varrho_{\kappa_k-1}^{(k)} = 1 \quad (\varrho_h^{(k)} = 1)$$

oder die Werte

$$(R_{II}) \quad \varrho_1^{(k)} = 2, \varrho_2^{(k)} = 1, \dots, \varrho_{\kappa_k-2}^{(k)} = 1, \varrho_{\kappa_k-1}^{(k)} = 2. \quad \begin{pmatrix} \varrho_{2^i-1}^{(k)} = 2 \\ \varrho_{2^i}^{(k)} = 1 \end{pmatrix}$$

Der Fall eines Restes (R_{II}) kann nur bei geradem κ_k eintreten. Nehmen wir an, von den λ Größen κ_k seien im ganzen λ_0 gerade, nämlich die folgenden

$$\kappa_{\tau_1}, \kappa_{\tau_2}, \dots, \kappa_{\tau_{\lambda_0}}.$$

Dann werden alle $\lambda - \lambda_0$ Formen $\Phi_{(k)}$, für welche der Index k keiner der Zahlen $\tau_1, \tau_2, \dots, \tau_{\lambda_0}$ gleich ist, gewiß Reste von der Gestalt (R_I) besitzen.

Für die in bezug auf 2 primitive Form f mögen die λ_f Formen

$$\Phi_{(\tau_{\pi_1})}, \Phi_{(\tau_{\pi_2})}, \dots, \Phi_{(\tau_{\pi_{\lambda_f}})}$$

Reste von der Art (R_{II}) lassen, während alle $\lambda - \lambda_f$ übrigen Formen $\Phi_{(k)}$ Reste von der Art (R_I) besitzen mögen. Wir können dann sagen, der Form f gehöre die Kombination

$$[\pi_1, \pi_2, \dots, \pi_{\lambda_f}]$$

an, und es ist sofort einleuchtend, daß durch diese Kombination von Zahlen π aus der Reihe der Zahlen $1, 2, \dots, \lambda_0$ die $n - 1$ Invarianten σ der Form f vollständig bestimmt sind. Wir erschließen daraus leicht den Satz:

E. Wenn die $n - 1$ Invarianten o_h und der Index I gegeben sind, so führen höchstens

$$2^{\lambda_0} \left(\leq 2^{\left[\frac{n}{2} \right]} \right)^*)$$

von den sämtlichen 2^{n-1} verschiedenen Kombinationen

$$(\sigma_1, \sigma_2, \dots, \sigma_{n-2}, \sigma_{n-1}); \quad \sigma_h = 1, 2,$$

welche sich überhaupt bilden lassen, zu wirklich existierenden Ordnungen

$$O: \left(\sigma_1, \sigma_2, \dots, \sigma_{n-1} \right), I.$$

In der Tat ist die Anzahl aller überhaupt zulässigen Kombinationen (σ_h) höchstens so groß wie die Anzahl sämtlicher angebbaren Kombinationen

$$[\pi_1, \pi_2, \dots, \pi_{\lambda_f}].$$

Diese letztere Anzahl ist aber offenbar gleich

$$\sum_{\lambda_f=1}^{\lambda_0} \frac{\lambda_0(\lambda_0-1) \cdots (\lambda_0-\lambda_f+1)}{1 \cdot 2 \cdots \lambda_f} = 2^{\lambda_0}.$$

Hieraus geht das Behauptete sofort hervor.

Wie wir später (im Kap. XI) nachweisen werden, ist die Anzahl derjenigen Kombinationen (σ_h), welche, mit den $n - 1$ Größen o_h und der Zahl I verbunden, zu wirklich existierenden Ordnungen führen, in der Tat mit nur wenigen Ausnahmen gleich 2^{λ_0} .

Eine dieser Ausnahmen kann eintreten, wenn $\lambda = \lambda_0$ ist. Alsdann sind sämtliche Zahlen $\pi_1, \pi_2, \dots, \pi_{\lambda}$ und mithin auch ihre Summe $\pi_1 + \pi_2 + \cdots + \pi_{\lambda} = n$ gerade, und es fragt sich, ob die Kombination $[1, 2, \dots, \lambda]$ eine wirklich bestehende Ordnung zu liefern vermag. Für diese Kombination müßte

$$\sigma_1 = 2, \sigma_2 = 1, \dots, \sigma_{n-2} = 1, \sigma_{n-1} = 2$$

sein, und es müßten sämtliche Reste $\Phi_{(k)}$ die Gestalt (R_{II}) erhalten. Für die Determinanten $|\Phi_{(k)}|$ dieser Formen werden daher die Kongruenzen gelten

$$|\Phi_{(k)}| \equiv (-1)^{\frac{\pi_k}{2}} \pmod{4},$$

aus welchen sich sofort

$$|\Phi_{(1)}| |\Phi_{(2)}| \cdots |\Phi_{(\lambda)}| \equiv (-1)^{\frac{\pi_1 + \pi_2 + \cdots + \pi_{\lambda}}{2}} \equiv (-1)^{\frac{n}{2}} \pmod{4}$$

*) Es ist $n \geq \pi_{\pi_1} + \pi_{\pi_2} + \cdots + \pi_{\pi_{\lambda_0}}$, und wir haben $\pi_{\pi_{\pi}} > 0$, $\pi_{\pi_{\pi}} \equiv 0 \pmod{2}$, also $\pi_{\pi_{\pi}} \geq 2$. Demnach kommt $n \geq 2\lambda_0$ oder $\lambda_0 \leq \left[\frac{n}{2} \right]$.

ergibt. Nun ist die Determinante der Form φ , wie man leicht erkennt,

$$(-1)^I \cdot d_{n-1} \equiv |\Phi_{(1)}| |\Phi_{(2)}| \cdots |\Phi_{(\lambda)}| \cdot 2^{\partial_{n-1}} \pmod{2^{2+\partial_{n-1}}} \\ [t > 1 + \partial_{n-1}].$$

Wir erhalten hieraus die Kongruenz

$$(-1)^I \cdot \frac{d_{n-1}}{2^{\partial_{n-1}}} \equiv (-1)^{\frac{n}{2}} \pmod{4}.$$

Beachten wir jetzt die Beziehungen

$$\frac{d_{n-1}}{2^{\partial_{n-1}}} = \left[\frac{o_1}{2^{\omega_1}} \right]^{n-1} \left[\frac{o_2}{2^{\omega_2}} \right]^{n-2} \cdots \left[\frac{o_{n-1}}{2^{\omega_{n-1}}} \right], \quad \left[\frac{o_h}{2^{\omega_h}} \right]^2 \equiv 1 \pmod{4}, \\ n \equiv 0 \pmod{2},$$

so gelangen wir zu der Bedingung

$$\prod_{h=1}^{\frac{n}{2}} \frac{o_{2h-1}}{2^{\omega_{2h-1}}} \equiv (-1)^{I + \frac{n}{2}} \pmod{4}.$$

Dieselbe liefert den folgenden Ausnahmefall:

Sobald $\lambda = \lambda_0$ [$n \equiv 0 \pmod{2}$] und

$$\prod_{h=1}^{\frac{n}{2}} \frac{o_{2h-1}}{2^{\omega_{2h-1}}} \equiv -(-1)^{I + \frac{n}{2}} \pmod{4}$$

wird, führen höchstens 2^{λ_0-1} Kombinationen der σ_h zu wirklich existierenden Ordnungen, da in diesem Falle die Kombination

$$\sigma_1 = 2, \sigma_2 = 1, \dots, \sigma_{n-2} = 1, \sigma_{n-1} = 2$$

unmöglich ist.

II. Den gewonnenen Resultaten wollen wir jetzt einen Ausdruck geben, welcher uns ohne weiteres in den Stand setzt, zu entscheiden, ob irgendeine Ordnung O Formen enthalten kann.

Eine Ordnung

$$\left(\begin{array}{l} \sigma_0 = 1; \sigma_1, \sigma_2, \dots, \sigma_{n-2}, \sigma_{n-1}; \sigma_n = 1 \\ o_0 = 0; o_1, o_2, \dots, o_{n-2}, o_{n-1}; o_n = 0 \end{array} \right), \quad I$$

ist nur dann möglich, wenn die ganzen Zahlen σ_h , o_h und I den folgenden Gesetzen gehorchen:

1) Die Zahlen σ_h und $\sigma_{h-1} o_h \sigma_{h+1}$ sind nicht beide zugleich durch 2 teilbar (σ_h relativ prim zu $\sigma_{h-1} o_h \sigma_{h+1}$).

2) Die Größen $\frac{\sigma_{h-1} o_h}{\sigma_{h+1}}$ und $\frac{\sigma_{h+1} o_h}{\sigma_{h-1}}$ sind ganze Zahlen.

3) Wenn $n \equiv 0 \pmod{2}$ und $\sigma_1 = 2, \sigma_2 = 1, \dots, \sigma_{n-2} = 1, \sigma_{n-1} = 2$ ist, so wird

$$\prod_{h=1}^{\frac{n}{2}} \frac{o_{2h-1}}{2^{\omega_{2h-1}}} \equiv (-1)^{I + \frac{n}{2}} \pmod{4}.$$

Damit diese Gesetze ohne Ausnahme auch für $h = 1$ und $h = n - 1$ gelten, haben wir den $2(n - 1)$ Invarianten o_h und σ_h noch je zwei weitere Invarianten $o_0 = 0, o_n = 0$ und $\sigma_0 = 1, \sigma_n = 1$ hinzugefügt.

Der Satz 1) zerfällt in zwei Teile:

α) Wenn $o_h \equiv 0 \pmod{2}$ ist, so wird $\sigma_h = 1$, und wenn $\sigma_h = 2$ ist, so wird $o_h \equiv 1 \pmod{2}$.

Es ergibt sich dieses aus der Relation $\sigma_{\vartheta_k} = 1$; denn sobald $o_h \equiv 0 \pmod{2}$ ist, muß h mit einer der Zahlen ϑ_k übereinstimmen, und sobald $\sigma_h = 2$ ist, muß h von den Zahlen ϑ_k verschieden und folglich $o_h \equiv 1 \pmod{2}$ sein.

β) Wenn $\sigma_h = 2$ ist, so wird $\sigma_{h-1} = 1, \sigma_{h+1} = 1$.

Dies erkennt man sofort aus dem Satze D. und den Werten, welche die Größen $\varrho_h^{(k)}$ annehmen können.

Den Satz 2) können wir auch folgendermaßen aussprechen:

Es ist $\sigma_{h-1} o_h \equiv \sigma_{h-1} o_h \sigma_{h+1} \equiv o_h \sigma_{h+1} \pmod{2}$;

oder:

Wenn $\sigma_{h-1} \neq \sigma_{h+1}$, also $\sigma_{h-1} \sigma_{h+1} = 2$ ist, so muß $o_h \equiv 0 \pmod{2}$ sein;

oder:

Wenn $\sigma_{h-1} o_h \sigma_{h+1}$ durch 2 und nicht durch 4 teilbar ist, so wird $\sigma_{h-1} = 1, \sigma_{h+1} = 1$.

Die Werte der $\varrho_h^{(k)}$ zeigen, daß stets $\sigma_{h-1} = \sigma_{h+1}$ ist, sobald h keiner der Zahlen ϑ_k gleich ist; es kann demnach nur dann $\sigma_{h-1} \neq \sigma_{h+1}$ sein, wenn h mit einer der Zahlen ϑ_k übereinstimmt, also $o_h \equiv 0 \pmod{2}$ ist.

Aus den Gesetzen 1) und 2) läßt sich der Satz E. von neuem herleiten.

Kap. V. Sätze über Hauptreste. — Grundformen für einen Modul N .

Die Hauptreste einer primitiven Formenklasse f besitzen eine Reihe interessanter Eigenschaften, von denen wir einige in dem Folgenden mitteilen wollen.

I. Wir setzen die höchste in dem Produkt $\sigma_{n-1} o_1 o_2 \dots o_{n-1}$ aufgehende Potenz einer Primzahl q gleich $q^{G(q)}$; für ein ungerades $q = p$ wird diese Potenz offenbar gleich $p^{v_{n-1}(p)}$, dagegen für $q = 2$ gleich $\sigma_{n-1} \cdot 2^{v_{n-1}(2)}$. Je nachdem eine Zahl $q^t \leq q^{G(q)}$ oder $> q^{G(q)}$ ist, treten in den Hauptresten von f für den Modul q^t mittlere Koeffizienten auf, welche

durch den Modul q^t teilbar sind, oder es ist keiner dieser Koeffizienten durch q^t teilbar. Aus diesem Grunde beschränken wir uns auf die Untersuchung von Moduln N , deren Primfaktoren $q^t (t > 0)$ die Potenzen $q^{G(q)}$ überschreiten.

Es möge die Form φ einen Hauptrepräsentanten der primitiven Klasse f in bezug auf einen derartigen Modul N bedeuten. Wir bezeichnen die aus den ersten $h (= 1, 2, \dots, n)$ Horizontal- und Vertikalreihen der Form φ gebildeten symmetrischen Unterdeterminanten durch $\sigma_h d_{h-1} \varphi_h$, und wir setzen noch $\varphi_0 = 1$. Die Zahl φ_n ist gleich $(-1)^I$.

Die sämtlichen $n + 1$ Größen $\varphi_h (h = 0; 1, 2, \dots, n)$ sind offenbar ganze Zahlen, und wir behaupten:

1. Die $n + 1$ Zahlen

$$\varphi_0; \varphi_1, \varphi_2, \dots, \varphi_{n-1}, \varphi_n$$

fallen sämtlich zu N relativ prim aus.

Zunächst möge N aus einer einzigen Primzahlpotenz $q^t (> q^{G(q)})$ bestehen. Ein Blick auf die Hauptrepräsentanten φ zeigt uns, daß die Determinanten $\sigma_h d_{h-1} \varphi_h$ sich für ein $q^t = p^t$ als Produkte aus den Potenzen $p^{2h-1(p)}$ und aus zu p primen Zahlen darstellen, während sie für ein $q^t = 2^t$ gleich Produkten aus den Potenzen $\sigma_h \cdot 2^{2h-1(2)}$ und aus ungeraden Zahlen werden. Hieraus schließen wir, daß die φ_h in der Tat zu der Primzahl q relativ prim sind. Falls N sich aber aus mehreren Primzahlpotenzen $q^t (> q^{G(q)})$ zusammensetzt, so ist der Hauptrepräsentant φ für den Modul N zugleich Hauptrepräsentant für jeden dieser Moduln q^t . Nach dem soeben Gesagten sind daher die Größen φ_h zu jeder der Primzahlen q prim, und sie werden demnach auch zu der Zahl N relativ prim sein.

Die n mittleren Koeffizienten einer beliebigen quadratischen Form können stets durch die $\frac{n(n-1)}{2}$ seitlichen Koeffizienten dieser Form und durch die aus den ersten $h (= 1, 2, \dots, n)$ Reihen dieser Form gebildeten symmetrischen Unterdeterminanten ausgedrückt werden. Beachten wir die besonderen Werte, welche die $\frac{n(n-1)}{2}$ seitlichen Koeffizienten in einem Hauptrest für den Modul N besitzen, so gelangen wir zu dem folgenden Satze:

2. Wenn die Form φ einen Hauptrepräsentanten der Klasse f für den Modul N vorstellt, so kann vermittels der $n + 1$ zur Form φ gehörigen Zahlen $\varphi_h (h = 0; 1, 2, \dots, n)$ ein Hauptrest der Klasse f für den Modul N gefunden werden.

Ist N beispielsweise ungerade, so liefert ein Hauptrepräsentant $\varphi \pmod{N}$ uns den Hauptrest

$$\varphi \equiv \begin{pmatrix} \sigma_1 \varphi_1 \\ \sigma_0 \varphi_0, \\ o_1 \cdot \frac{\sigma_2 \varphi_2}{\sigma_1 \varphi_1}, \\ o_1 o_2 \cdot \frac{\sigma_3 \varphi_3}{\sigma_2 \varphi_2}, \\ \dots \dots, \\ o_1 o_2 \dots o_{n-1} \cdot \frac{\sigma_n \varphi_n}{\sigma_{n-1} \varphi_{n-1}} \end{pmatrix} \pmod{N}.$$

Ein wenig umständlicher gestaltet sich der Ausdruck des Hauptrestes $(\text{mod } N)$ durch die Größen φ_h , wenn der Modul N gerade ist und nicht alle Invarianten σ_h der Form φ gleich 1 werden. Wenn $N = 2^t$ ist, müssen wir, um vermittle der Zahlen φ_h eines Hauptrepräsentanten einen Hauptrest zu finden, so viele willkürliche ungerade Zahlen $\mathfrak{A}$ einführen, als es Invarianten σ_h gibt, welche gleich 2 ausfallen.

Wir nennen eine Form $\psi (\sim f)$ eine *Grundform* der Klasse f für einen Modul N , sobald die aus den ersten h Horizontal- und Vertikalreihen der Form ψ gebildeten symmetrischen Determinanten Werte $\sigma_h d_{h-1} \psi_h$ erhalten, in welchen die Zahlen ψ_h zu N relativ prim sind.

Es ist leicht einzusehen, daß die Klasse f für einen jeden Modul N Grundformen enthält; denn treten in dem Modul N die Primzahlen $q_1, q_2, \dots, q_{c_0}$ auf, so bildet nach dem Satz 1. offenbar jeder Hauptrepräsentant in bezug auf einen Modul $N^* = q_1^{t_1} q_2^{t_2} \dots q_{c_0}^{t_{c_0}} [t_c > G(q_c), c = 1, 2, \dots, c_0]$ eine Grundform für den Modul N .

Wir haben die folgenden Sätze: Eine Grundform für einen Modul N ist Grundform für alle in N enthaltenen Faktoren; und umgekehrt: Eine Form, welche für gewisse Primzahlen q eine Grundform ist, stellt auch für jeden Modul N , der nur durch diese Primzahlen q teilbar ist, eine Grundform vor.

Mit Hilfe von Betrachtungen, welche denen in Kap. II angestellten analog sind, kann man den Satz beweisen:

3. Jede Grundform ψ für einen Modul N kann vermittle einer Substitution

$$S = \begin{pmatrix} 1, s_1^2, s_1^3, \dots, s_1^n \\ 1, s_2^3, \dots, s_2^n \\ 1, \dots, s_3^n \\ \dots \dots \dots \\ 1 \end{pmatrix}$$

in einen Hauptrepräsentanten φ für einen jeden Modul N_0 , der keine anderen Primzahlen enthält als der Modul N , übergeführt werden.

Man sieht, daß die Substitution S die Werte der $n + 1$ Größen ψ_h

ungeändert läßt, so daß für die Form $\varphi = \bar{S}\psi S$ die Relationen $\varphi_h = \psi_h$ gelten. Da nun einerseits nach 2. aus den Größen φ_h eines Hauptrepräsentanten $\varphi \pmod{N_0}$ sofort ein Hauptrest der Klasse f für den Modul N_0 gefunden werden kann und andererseits die Zahlen ψ_h der Grundform $\psi \pmod{N}$ mit den Zahlen φ_h übereinstimmen, so können wir auch unmittelbar von der Grundform ψ ohne Zuhilfenahme der Form φ zu Hauptresten der Klasse f für den Modul N_0 gelangen, und wir erhalten den Satz:

4. Eine jede Grundform $\psi \pmod{N}$ führt unmittelbar zu Hauptresten in bezug auf jeden Modul N_0 , welcher keine anderen Primzahlen enthält als der Modul N .

Insbesondere schließen wir hieraus:

F. Eine jede Grundform ψ für eine Primzahl q liefert Hauptreste für alle Moduln q^t .

Die Bestimmung einer Grundform ψ für eine Primzahl q kann sehr leicht geschehen. Denn es besteht der folgende Satz, in welchem $\bar{\lambda}(q) - 1$ die Zahl der durch q teilbaren Größen $\sigma_{h-1} o_h \sigma_{h+1}$ ($h = 1, 2, \dots, n-1$) bezeichnet:

Eine jede primitive Form f kann durch höchstens $n - \bar{\lambda}(q)$ Substitutionen von der Art

$$S_{(i,k)}^{\pm 1}: \quad x_i = x'_i, \quad x_k = \pm x'_i + x'_k; \quad x_h = x'_h \quad (h \neq i, k)$$

und durch höchstens $n - 1$ Substitutionen von der Art

$$P_{(l,m)}: \quad x_l = x'_m, \quad x_m = -x'_l; \quad x_h = x'_h \quad (h \neq l, m)$$

in eine Grundform ψ für die Primzahl q transformiert werden.

Die Substitutionen $P_{(l,m)}$ sind offenbar sehr einfach auszuführen, da sie im Grunde nur eine Vertauschung zweier Reihen der quadratischen Koeffizientensysteme bedeuten.

II. Wenn für eine primitive Form f die Zahl $\sigma_{h-1} o_h \sigma_{h+1}$ durch eine Primzahl q teilbar ist, so gibt es unter den h -reihigen symmetrischen Determinanten $D_h(f)$ der Form f solche, für welche die Zahl $\frac{D_h(f)}{\sigma_h d_{h-1}} = \Delta_h(f)$ zu q relativ prim wird.

Beweis: Es sei zunächst $q = p$. Ein Hauptrepräsentant $\varphi \pmod{p^t > p^{G(p)}}$ der Klasse f läßt einen Rest

$$\varphi \equiv \begin{pmatrix} \alpha_1, \\ \cdot \quad \cdot \quad \cdot, \\ p^{\omega_1 + \omega_2 + \dots + \omega_{h-1}} \alpha_h, \\ \cdot \quad \cdot \quad \cdot, \\ p^{\omega_1 + \omega_2 + \dots + \omega_h} \alpha_{h+1}, \\ \cdot \quad \cdot \quad \cdot, \\ p^{\omega_1 + \omega_2 + \dots + \omega_{n-1}} \alpha_n \end{pmatrix} \pmod{p^t}. \quad (3^*)$$

Aus demselben ersehen wir, daß die höchste Potenz von p , welche in der Determinante

$$\varphi \begin{pmatrix} 1, 2, \dots, h \\ 1, 2, \dots, h \end{pmatrix} = \sigma_h d_{h-1} \varphi_h$$

aufgeht, gleich $p^{\partial h-1}$ ist, während die höchste, in allen übrigen h -reihigen Determinanten

$$\varphi \begin{pmatrix} i_1, i_2, \dots, i_h \\ k_1, k_2, \dots, k_h \end{pmatrix}$$

der Form φ enthaltene Potenz von p gleich $p^{\partial h-1+\omega_h}$ wird.

Geht φ in die Form f durch die Substitution

$$\xi_i = \sum_{k=1}^n u_i^k x_k$$

über, so bekommen wir für die symmetrischen h -reihigen Unterdeterminanten der Form f die Gleichungen

$$D(l_1, l_2, \dots, l_h) = \frac{\sum_{(i,k)}^{1,n} \varphi \begin{pmatrix} i_1, i_2, \dots, i_h \\ k_1, k_2, \dots, k_h \end{pmatrix} U \begin{pmatrix} l_1, l_2, \dots, l_h \\ i_1, i_2, \dots, i_h \end{pmatrix} U \begin{pmatrix} l_1, l_2, \dots, l_h \\ k_1, k_2, \dots, k_h \end{pmatrix}}{(1 \cdot 2 \dots h) \cdot (1 \cdot 2 \dots h)}.$$

Dieselben liefern infolge der soeben gemachten Bemerkung die Kongruenz

$$\begin{aligned} D(l_1, l_2, \dots, l_h) &= \sigma_h d_{h-1} \Delta(l_1, l_2, \dots, l_h) \\ &\equiv \sigma_h d_{h-1} \varphi_h \cdot \left(U \begin{pmatrix} l_1, l_2, \dots, l_h \\ 1, 2, \dots, h \end{pmatrix} \right)^2 \pmod{p^{\partial h-1+\omega_h}} \end{aligned}$$

oder

$$(2) \quad \Delta(l_1, l_2, \dots, l_h) \equiv \varphi_h \cdot \left(U \begin{pmatrix} l_1, l_2, \dots, l_h \\ 1, 2, \dots, h \end{pmatrix} \right)^2 \pmod{p^{\omega_h(p)}}.$$

Die Zahlen $U \begin{pmatrix} l_1, l_2, \dots, l_h \\ 1, 2, \dots, h \end{pmatrix}$ können keinen gemeinsamen, von 1 verschiedenen Teiler besitzen; denn ein solcher Teiler müßte zugleich die Determinante $|u_i^k|$ teilen, und diese könnte mithin nicht gleich 1 sein.

Es müssen sich demnach unter den Zahlen $U \begin{pmatrix} l_1, l_2, \dots, l_h \\ 1, 2, \dots, h \end{pmatrix}$ solche befinden, welche zu p prim ausfallen. Wenn nun die Größe $\sigma_{h-1} o_h \sigma_{h+1}$ durch p teilbar ist, so wird offenbar $\omega_h(p) \geq 1$, und es entspricht alsdann infolge der Kongruenz (2) einer jeden durch p nicht teilbaren Zahl $U \begin{pmatrix} l_1, l_2, \dots, l_h \\ 1, 2, \dots, h \end{pmatrix}$ eine durch p nicht teilbare Zahl $\Delta(l_1, l_2, \dots, l_h)$. Hieraus erhellt die Richtigkeit unseres Satzes für den Fall $q = p$.

Wenn $q = 2$ ist, so gelangen wir durch Betrachtung eines Hauptrepräsentanten $\varphi \pmod{2^e > 2^{G(2)}}$ in ähnlicher Weise zu einer Kongruenz

$$(3) \quad \Delta(l_1, l_2, \dots, l_h) \equiv \varphi_h \cdot X_h^2 \pmod{\sigma_{h-1} 2^{\omega_h(2)} \sigma_{h+1}},$$

und diese Kongruenz zeigt uns, daß unter den Größen $\Delta(l_1, l_2, \dots, l_h)$ sich ungerade befinden müssen, sobald $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{2}$ wird. — Dieses letzte Resultat erhalten wir einfacher, wenn wir beachten, daß einerseits σ_h gleich 2 sein müßte, sobald alle $\Delta(l_1, l_2, \dots, l_h)$ gerade werden, während andererseits nach Kap. IV, Absatz II zu einem $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{2}$ stets ein σ_h gleich 1 gehört.

Aus jeder Grundform ψ in bezug auf eine Primzahl q können wir nach dem Satz F. für einen jeden Modul q^t Hauptrepräsentanten φ herleiten, deren Zahlen φ_h den Zahlen ψ_h gleich werden. Infolge dieses Umstandes dürfen wir die Kongruenzen (2) und (3) durch die Kongruenzen

$$(4) \quad \frac{D_h(f)}{\sigma_h d_{h-1}} = \Delta_h(f) \equiv \psi_h \cdot X_h^2 \pmod{p^{w_h(p)}}$$

und

$$(5) \quad \frac{D_h(f)}{\sigma_h d_{h-1}} = \Delta_h(f) \equiv \psi_h \cdot X_h^2 \pmod{\sigma_{h-1} 2^{w_h(2)} \sigma_{h+1}}$$

ersetzen, welche resp. für eine Grundform $\psi \pmod{p}$ und $\psi \pmod{2}$ gelten werden.

Für eine Grundform ψ in bezug auf den Modul $2 \prod_{h=1}^{n-1} o_h$ müssen offenbar sämtliche Kongruenzen (4) und (5) zu gleicher Zeit statthaben, und wir bekommen daher für eine Grundform $\psi \pmod{2 \prod_{h=1}^{n-1} o_h}$ die Beziehungen

$$\frac{D_h(f)}{\sigma_h d_{h-1}} = \Delta_h(f) \equiv \psi_h \cdot X_h^2 \pmod{\sigma_{h-1} o_h \sigma_{h+1}};$$

dieselben liefern uns für irgend zwei beliebige Zahlen $\Delta'_h(f)$ und $\Delta''_h(f)$ eine Kongruenz

$$\Delta'_h(f) \cdot \Delta''_h(f) \equiv [\Delta'_h(f)]^2 \pmod{\sigma_{h-1} o_h \sigma_{h+1}}.$$

Kap. VI. Mengengruppen für einen Modul N . — Charaktere.

I. Zwei Formen $\varphi = \{\alpha_{ik}\}$ und $\psi = \{\beta_{ik}\}$ heißen *kongruent* nach einem Modul N , wenn sie den sämtlichen Kongruenzen $\alpha_{ik} \equiv \beta_{ik} \pmod{N}$ genügen, d. h. wenn sie für den Modul N den nämlichen Rest ergeben. Gleicherweise nennen wir zwei *Formenklassen* f und g für einen Modul N *kongruent*, wenn sie irgend zwei nach diesem Modul N kongruente Repräsentanten enthalten. Wir drücken die Kongruenz zweier Formen φ und ψ durch das Zeichen $\equiv [\varphi \equiv \psi \pmod{N}]$ und die Kongruenz zweier Klassen f und g durch das Zeichen $\simeq [f \simeq g \pmod{N}]$ aus.

Wir beweisen den Satz:

Wenn für irgend zwei Repräsentanten φ und ψ zweier Klassen f und g eine Kongruenz

$$\varphi \equiv \psi \pmod{N}$$

besteht, so kann zu jeder Form f_0 der einen eine Form g_0 der anderen Klasse bestimmt werden, welche ihr nach dem Modul N kongruent ist.

Denn wenden wir irgendeine Substitution S gleichzeitig auf beide Formen φ und ψ an, so gehen dieselben in zwei Formen $\varphi_0 (\sim \varphi)$ und $\psi_0 (\sim \psi)$ über, welche ebenso wie φ und ψ einer Kongruenz

$$\varphi_0 \equiv \psi_0 \pmod{N}$$

genügen. Nun können wir S so wählen, daß $\varphi_0 = f_0$ wird. Dann wird das betreffende ψ_0 der Klasse g gleich $g_0 \equiv f_0 \pmod{N}$ werden.

Nach diesem Satze enthalten zwei Formenklassen, welche für einen Modul N kongruent sind, lauter gleiche Formenreste für den Modul N . Daraus folgt:

G. Sind zwei Formenklassen g , und $g_{\prime\prime}$ einer dritten Formenklasse f nach dem Modul N kongruent, so sind sie auch untereinander nach dem Modul N kongruent.

Denn ist $g, \simeq f \pmod{N}$ und $g_{\prime\prime}, \simeq f \pmod{N}$, so stimmen sowohl die Reste der Klasse g , für den Modul N als auch die Reste der Klasse $g_{\prime\prime}$, für den Modul N völlig mit den Resten der Klasse f für den Modul N überein. Es werden demnach auch die Reste der beiden Klassen g , und $g_{\prime\prime}$, für den Modul N einander gleich sein, d. h. es wird wirklich $g, \simeq g_{\prime\prime} \pmod{N}$ werden.

Der Satz G. zeigt uns die Möglichkeit einer Einteilung aller Formen in eine endliche Anzahl von *Gruppen* in bezug auf einen Modul N , indem wir zwei Formenklassen in dieselbe oder in verschiedene Gruppen $\pmod{N}$ aufnehmen, je nachdem sie in bezug auf N kongruent sind oder nicht. Es werden alsdann sämtliche Klassen einer und derselben Gruppe modulo N untereinander kongruent sein, während zwei Klassen aus verschiedenen Gruppen modulo N inkongruent sein werden.

Es leuchtet ein, daß die Anzahl aller verschiedenen Gruppen für einen Modul N , zu welchen wir durch diese Einteilung gelangen, notwendig eine endliche ist. Denn es gibt gewiß nicht mehr als $N^{\frac{n(n+1)}{2}}$ verschiedene Gruppen für den Modul N , da doch eine jede quadratische Form $f = \{a_{ik}\}$ nach dem Modul N mindestens einer der $N^{\frac{n(n+1)}{2}}$ Formen $\{R_{ik}\}$ ($R_{ik} = 1, 2, \dots, N$) kongruent sein muß.

II. Wenn der Modul N ein Produkt aus c_0 zueinander relativ primen Zahlen $N_1, N_2, \dots, N_{c_0}$ ist, so erfordert das Bestehen einer Kongruenz

$$f \simeq g \pmod{N}$$

das gleichzeitige Bestehen der c_0 einzelnen Kongruenzen

$$f \simeq g \pmod{N_c} \quad (c = 1, 2, \dots, c_0),$$

und umgekehrt findet die erstere Kongruenz wirklich statt, sobald die c_0 letzteren Kongruenzen sämtlich erfüllt sind.

In der Tat ist jeder Rest der Formenklassen f und g in bezug auf den Modul N zugleich Rest dieser Formenklasse für die sämtlichen Moduln N_c . Ist also $f \simeq g \pmod{N}$, so wird auch $f \simeq g \pmod{N_c}$ sein.

Falls umgekehrt die c_0 Kongruenzen

$$f \simeq g \pmod{N_c} \quad (c = 1, 2, \dots, c_0)$$

alle erfüllt sind, so können wir zunächst in der Klasse f gewiß c_0 Formen φ_c und in der Klasse g ebenfalls c_0 Formen ψ_c angeben derart, daß die c_0 Kongruenzen

$$\varphi_c \equiv \psi_c \pmod{N_c} \quad (c = 1, 2, \dots, c_0)$$

statthaben. Darauf können wir nach Kap. III, Absatz I in den Klassen f und g je eine Form φ und ψ angeben, welche den Kongruenzen

$$\varphi \equiv \varphi_c \pmod{N_c}, \quad \psi \equiv \psi_c \pmod{N_c} \quad (c = 1, 2, \dots, c_0)$$

genügen. Die Koeffizienten dieser beiden letzten Formen φ und ψ sind nun in bezug auf die sämtlichen zueinander primen Moduln N_c , also auch in bezug auf den Modul $N = N_1 N_2 \dots N_{c_0}$ kongruent; mithin wird

$$\varphi \equiv \psi \pmod{N},$$

und hieraus ergibt sich auf der Stelle

$$f \simeq g \pmod{N}.$$

Wählen wir für die Zahlen N_c die verschiedenen in N enthaltenen Primzahlpotenzen q^t , so zeigt sich, daß die Untersuchung von Gruppen für einen beliebigen Modul N auf die Untersuchung von Gruppen in bezug auf Primzahlpotenzen q^t hinausläuft.

III. Wir wollen die notwendigen und hinreichenden Bedingungen für die Kongruenz zweier Formenklassen nach einem Modul N auffinden. Nach Aufstellung dieser Bedingungen werden wir dazu übergehen, zu prüfen, ob es verschiedene Formenklassen von demselben Index I gibt, welche nach allen überhaupt möglichen Moduln kongruent sind, und aus diesen nach jedem Modul kongruenten Formenklassen werden wir die *Genera* von Formen bilden.

Nehmen wir an, die Klassen f und g seien nach dem Modul N kongruent, und setzen wir ferner, der größeren Einfachheit halber, voraus, daß f und g in bezug auf N primitiv sind und daß die Faktoren q^t ($t > 0$) von N für ein $q = p$ die Potenzen $p^{e_n-1(p)}$ und für ein $q = 2$ die Potenzen $\frac{4}{\sigma_{n-1}} 2^{e_n-1(2)}$, welche den Formen f und g entsprechen, übersteigen.

1. Nach den soeben bewiesenen Sätzen besitzen die Formenklassen f und g für einen jeden der in N aufgehenden Moduln q^t vollständig dieselben Reste und folglich auch dieselben Hauptreste. Aus der Gestalt dieser Hauptreste für die Moduln $q^t (> q^{G(2)})$ können wir aber die Werte der Zahlen

$$q^{\omega_1}, q^{\omega_2}, \dots, q^{\omega_{n-1}}$$

und, wenn $q = 2$ ist, auch die Werte der Zahlen

$$\sigma_1, \sigma_2, \dots, \sigma_{n-1},$$

welche zu den beiden Klassen f und g gehören, unmittelbar ablesen. Folglich müssen die Formen f und g für alle in N aufgehenden Primzahlen q dieselben Größen q^{ω_h} und, falls $q = 2$ ist, auch dieselben Größen σ_h darbieten.

2. Es seien $\Delta(f)$ und $\Delta(g)$ die Determinanten der Formen f und g . Für $\Delta(g) = |b_{ik}|$ besteht die Beziehung

$$\Delta(g) = \sum_{i=1}^n b_{ii} \frac{\partial \Delta(g)}{\partial b_{ii}} + 2 \sum_{i < k}^{1, n} b_{ik} \frac{\partial \Delta(g)}{\partial b_{ik}}.$$

Hierin ist die größte Potenz einer Primzahl q , welche in allen $\frac{\partial \Delta(g)}{\partial b_{ii}}$, $2 \frac{\partial \Delta(g)}{\partial b_{ik}}$ aufgeht, für ein $q = p$ gleich $p^{\partial_{n-2}(p)}$ und für ein $q = 2$ gleich $\sigma_{n-1} \cdot 2^{\partial_{n-2}(2)}$. Verändern wir also die Koeffizienten der Form g um Vielfache der Größe q^t , so wird die Determinante von g sich um Vielfache der Größe $p^{t + \partial_{n-2}(p)}$ oder der Größe $\sigma_{n-1} \cdot 2^{t + \partial_{n-2}(2)}$ ändern, je nachdem $q = p$ oder $q = 2$ ist. Erinnern wir uns, daß es zu f äquivalente Formen f_0 gibt, für welche $f_0 \equiv g \pmod{N}$ ist, so schließen wir daraus im Falle $q = p$ die Kongruenz

$$\Delta(f) \equiv \Delta(g) \pmod{p^{t + \partial_{n-2}(p)}}$$

und im Falle $q = 2$ die Kongruenz

$$\Delta(f) \equiv \Delta(g) \pmod{\sigma_{n-1} \cdot 2^{t + \partial_{n-2}(2)}}.$$

3. Man kann gewisse als Funktionen der Koeffizienten der Formen f und g ausdrückbare Einheiten angeben, welche für f und für g dieselben Werte annehmen.

Wir bezeichnen durch $\sigma_h d_{h-1} \Delta_h(f)$ die symmetrischen h -reihigen Unterdeterminanten, welche sich aus dem quadratischen System der Form f bilden lassen, und durch $\sigma_h d_{h-1}(f_h)$ alle aus h Reihen einer beliebigen Form der Klasse f bestehenden Unterdeterminanten. Die Reste der Zahlen $\sigma_h d_{h-1}(f_h)$ für einen Modul N sind offenbar durch die Reste der Formenklasse f für den Modul N schon völlig bestimmt. Wenn also zwei Formenklassen f und g für den Modul N gleiche Reste lassen, so nehmen auch die Zahlen $\sigma_h d_{h-1}(f_h)$ und $\sigma_h d_{h-1}(g_h)$ für den Modul N

kongruente Werte an. Daher sind die Eigenschaften der Kongruenzen

$$\sigma_h d_{h-1}(f_h) \equiv \sigma_h d_{h-1} m_h \pmod{N}$$

für die Gruppe, welcher die Form f für den Modul N angehört, charakteristisch.

Wir sprechen den Satz aus:

Wenn die Kongruenz $(f_h) \equiv m_h \pmod{N_0}$ [$1 \leq h \leq n-1$] für ein bestimmtes m_h Lösungen besitzt, so sind auch alle weiteren Kongruenzen $(f_h) \equiv m_h Z^2 \pmod{N_0}$ lösbar, in welchen Z eine beliebige zu N_0 relativ prime Größe bedeutet.

Denn besitzt die Kongruenz $(f_h) \equiv m_h \pmod{N_0}$ Lösungen, so muß es in der Klasse f eine Form Φ geben, welche eine h -reihige symmetrische Determinante $\sigma_h d_{h-1} \Delta_h(\Phi) \equiv \sigma_h d_{h-1} m_h \pmod{\sigma_h d_{h-1} N_0}$ aufweist. Da die Zahl $h \geq 1$ und $\leq n-1$ ist, können wir annehmen, daß in dieser Determinante Glieder aus einer gewissen, etwa der i^{ten} Reihe des Systems Φ , dagegen keine Glieder einer gewissen andern, etwa der k^{ten} Reihe dieses Systems erscheinen. Weil die Zahl Z zu N_0 relativ prim sein soll, können wir eine Zahl Z_0 bestimmen, welche der Kongruenz

$$ZZ_0 \equiv 1 \pmod{N_0^2}$$

oder einer Gleichung

$$ZZ_0 - z N_0 z_0 N_0 = 1$$

genügt. Durch die Substitution

$$x_i = Zx'_i + z N_0 x'_k, \quad x_k = z_0 N_0 x'_i + Z_0 x'_k \\ \left[x_i \equiv Zx'_i, \quad x_k \equiv \frac{1}{Z} x'_k \pmod{N_0} \right]$$

von der Determinante 1 geht dann die Form Φ in eine Form Φ_0 über, in welcher an die Stelle der Determinante $\sigma_h d_{h-1} \Delta_h(\Phi)$ offenbar eine Determinante

$$\sigma_h d_{h-1} \Delta_h(\Phi_0) \equiv \sigma_h d_{h-1} \Delta_h(\Phi) Z^2 \equiv \sigma_h d_{h-1} m_h Z^2 \pmod{\sigma_h d_{h-1} N_0}$$

getreten ist. Die Zahl $\Delta_h(\Phi_0)$ ist gleichfalls eine der Zahlen (f_h) , und es besitzt sonach die Kongruenz $(f_h) \equiv m_h Z^2 \pmod{N_0}$ in der Tat Lösungen.

Aus dem vorstehenden Satz erschen wir, daß bei einer Untersuchung der Kongruenzen $(f_h) \equiv m_h \pmod{N_0}$ vor allem der quadratische Charakter der Zahl m_h in Betracht kommen wird. Nehmen wir beispielsweise an, der Modul N_0 sei eine Potenz einer Primzahl p ($N_0 = p^t$) und die Kongruenz $(f_h) \equiv m_h \pmod{p^t}$ nicht für alle zu p primen Zahlen m_h lösbar. Dann ist diese Kongruenz entweder nur für alle solche zu p primen m_h lösbar, welche quadratische Reste von p sind, oder nur für alle solche zu p primen m_h , welche quadratische Nichtreste von p sind. Wenn wir den

Formen f im ersten Fall einen Charakter $\left(\frac{f_h}{p}\right) = +1$ zuteilen, dagegen im zweiten Fall einen Charakter $\left(\frac{f_h}{p}\right) = -1$, so kann die Einheit $\left(\frac{f_h}{p}\right)$, welche der Form f entspricht, dazu dienen, die Gruppe, welcher die Klasse f für den Modul $p^{e_0 + \partial_{h-1}(p)}$ angehört, von anderen Gruppen nach diesem Modul zu trennen.

Wenn wir zeigen sollen, daß die Größen (f_h) einer Klasse f gewisse gegebene Charaktere besitzen, so genügt es, die beiden folgenden Punkte festzustellen:

1. daß die Zahlen $\Delta_h(\varphi)$ einer bestimmten Form φ der Klasse f wirklich die gegebenen Charaktere besitzen;

2. daß, wenn die Zahlen $\Delta_h(\Phi)$ einer beliebigen Form Φ der Klasse f diese Charaktere besitzen, alsdann auch die Zahlen $\Delta_h(F)$ aller aus Φ vermittels Substitutionen

$$S_{(i,k)}^{\pm 1}: \quad x_i = x'_i, \quad x_k = \pm x'_i + x'_k$$

hervorgehenden Formen F dieselben Charaktere besitzen.

Sind diese beiden Punkte festgestellt, so müssen die sämtlichen (f_h) der Klasse f wirklich die gegebenen Charaktere besitzen. Man sieht dieses auf der Stelle ein, indem man von dem bekannten Satze Gebrauch macht, daß eine jede ganzzahlige Substitution von der Determinante 1 sich in eine Reihe von Substitutionen $S_{(i,k)}^{\pm 1}$ zerlegen läßt. — Diesen allgemeinen Weg zur Entscheidung über die Existenz von Charakteren schlagen wir etzt in einigen Beispielen ein.

Blicken wir zunächst auf die Hauptrepräsentanten φ für einen Modul q' zurück, so machen wir die folgenden Beobachtungen:

($q = p$). Wenn $o_h \equiv 0 \pmod{p}$ ist, so werden alle Zahlen $\Delta_h(\varphi)$ eines Hauptrepräsentanten $\varphi \pmod{p' > p^{e_{n-1}(p)}}$ mit Ausnahme einer einzigen kongruent Null $\pmod{p}$, und dieses eine zu p prime $\Delta_h(\varphi)$ besitzt selbstverständlich einen bestimmten Charakter $\left(\frac{\Delta_h(\varphi)}{p}\right)$.

($q = 2$). Wenn $\sigma_{h-1} o_h \sigma_{h+1} \equiv 4 \pmod{8}$ ist, so sind die $\Delta_h(\varphi)$ eines Hauptrepräsentanten $\varphi \pmod{2^{e' > \sigma_{n-1} \cdot 2^{e_{n-1}(2)}}$ zum Teil $\equiv 0 \pmod{4}$, zum Teil $\equiv \delta (= \pm 1) \pmod{4}$. Wenn $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{8}$ ist, so werden alle $\Delta_h(\varphi)$ eines Hauptrepräsentanten $\varphi \pmod{2^{e' > \sigma_{n-1} \cdot 2^{e_{n-1}(2)}}$ mit Ausnahme eines einzigen $\equiv 0 \pmod{4}$, und dieses eine ist ungerade. Im Falle $\sigma_{h-1} o_h \sigma_{h+1} \equiv 4 \pmod{8}$ besitzen also die ungeraden $\Delta_h(\varphi)$ einen Charakter $(-1)^{\frac{\Delta_h(\varphi)-1}{2}}$ und im Falle $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{8}$ zwei Charaktere $(-1)^{\frac{\Delta_h(\varphi)-1}{2}}$ und $\left(\frac{2}{\Delta_h(\varphi)}\right)$, während die geraden $\Delta_h(\varphi)$ in diesen beiden Fällen $\equiv 0 \pmod{4}$ sind.

Diese Bemerkungen, welche sich unmittelbar aus der Gestalt der Hauptrepräsentanten φ ergeben, lassen vermuten, daß einer Formenklasse f im Falle $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{p}$ ein Charakter $\left(\frac{f_h}{p}\right)$, im Falle $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{4}$ ein Charakter $(-1)^{\frac{(f_h)-1}{2}}$ und im Falle $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{8}$ ein Charakter $\left(\frac{2}{f_h}\right)$ angehört. Wir bestätigen diese Vermutung, indem wir den Satz beweisen:

Ist $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{p}$ oder $\equiv 0 \pmod{4}$ oder $\equiv 0 \pmod{8}$, so besitzen die Größen $\Delta_h(F)$ einer Form F , welche aus einer anderen Form $\Phi(\sim f)$ mittels einer Substitution $S_{(i,k)}^{\pm 1}$ hervorgeht, dieselben quadratischen Charaktere in bezug auf die Moduln p oder 4 oder 8 wie die Zahlen $\Delta_h(\Phi)$.

Beweis: Die h -reihigen symmetrischen Unterdeterminanten der Form F ,

$$\sigma_h d_{h-1} F_h = D_F \begin{pmatrix} k_0, k_1, \dots, k_{h-1} \\ k_0, k_1, \dots, k_{h-1} \end{pmatrix},$$

werden gleich den Größen

$$(6) \quad \sigma_h d_{h-1} \Phi_h = D_\Phi \begin{pmatrix} k_0, k_1, \dots, k_{h-1} \\ k_0, k_1, \dots, k_{h-1} \end{pmatrix},$$

sobald sich unter den Zahlen $k_0, k_1, \dots, k_{h-1}$ entweder sowohl i als k oder weder i noch k befinden; dagegen werden diese Determinanten gleich den Größen

$$(7) \quad D_\Phi \begin{pmatrix} i, k_1, \dots, k_{h-1} \\ i, k_1, \dots, k_{h-1} \end{pmatrix} \pm 2 D_\Phi \begin{pmatrix} i, k_1, \dots, k_{h-1} \\ k, k_1, \dots, k_{h-1} \end{pmatrix} + D_\Phi \begin{pmatrix} k, k_1, \dots, k_{h-1} \\ k, k_1, \dots, k_{h-1} \end{pmatrix} \\ = \sigma_h d_{h-1} \left(\Phi_h' \pm \frac{2}{\sigma_h} \Phi_h'' + \Phi_h'' \right),$$

sobald $k_0 = i$ wird. Es ist klar, daß unser Satz für diejenigen Zahlen F_h statthat, welche aus den ersteren Größen D_F entspringen, da sie mit gewissen der Zahlen Φ_h übereinstimmen. Die Zahlen F_h dagegen, welche aus den letzteren Größen D_F entspringen, gewinnen mit Hilfe des bekannten Determinantensatzes

$$D_\Phi \begin{pmatrix} i, k_1, \dots, k_{h-1} \\ i, k_1, \dots, k_{h-1} \end{pmatrix} D_\Phi \begin{pmatrix} k, k_1, \dots, k_{h-1} \\ k, k_1, \dots, k_{h-1} \end{pmatrix} - \left[D_\Phi \begin{pmatrix} i, k_1, \dots, k_{h-1} \\ k, k_1, \dots, k_{h-1} \end{pmatrix} \right]^2 \\ = D_\Phi \begin{pmatrix} k_1, \dots, k_{h-1} \\ k_1, \dots, k_{h-1} \end{pmatrix} D_\Phi \begin{pmatrix} i, k, k_1, \dots, k_{h-1} \\ i, k, k_1, \dots, k_{h-1} \end{pmatrix}$$

oder

$$(8) \quad \sigma_h^2 \Phi_h' \Phi_h'' - (\Phi_h'')^2 = \sigma_{h-1} o_h \sigma_{h+1} \cdot \Phi_{h-1} \Phi_{h+1},$$

die Gestalt

$$F_h = \Phi_h' \left(1 \pm \frac{\Phi_h'''}{\sigma_h \Phi_h'} \right)^2 + \frac{\sigma_{h-1} o_h \sigma_{h+1} \cdot \Phi_{h-1} \Phi_{h+1}}{\sigma_h^2 \Phi_h'},$$

$$F_h = \Phi_h'' \left(1 \pm \frac{\Phi_h'''}{\sigma_h \Phi_h''} \right)^2 + \frac{\sigma_{h-1} o_h \sigma_{h+1} \cdot \Phi_{h-1} \Phi_{h+1}}{\sigma_h^2 \Phi_h''}.$$

Ist jetzt zunächst $o_h \equiv 0 \pmod{p}$, so sind die Zahlen Φ_h' und Φ_h'' entweder beide $\equiv 0 \pmod{p}$, oder es wird wenigstens eine derselben relativ prim zu p . Im ersteren Falle wird infolge der Kongruenz

$$\sigma_h^2 \Phi_h' \Phi_h'' - (\Phi_h''')^2 \equiv 0 \pmod{o_h}$$

auch $\Phi_h''' \equiv 0 \pmod{p}$, und der Ausdruck $F_h = \Phi_h' \pm \frac{2}{\sigma_h} \Phi_h''' + \Phi_h''$ ist also gleichfalls $\equiv 0 \pmod{p}$; im letzteren Falle ist entweder

$$\frac{\sigma_{h-1} o_h \sigma_{h+1} \Phi_{h-1} \Phi_{h+1}}{\sigma_h^2 \Phi_h'} \quad \text{oder} \quad \frac{\sigma_{h-1} o_h \sigma_{h+1} \Phi_{h-1} \Phi_{h+1}}{\sigma_h^2 \Phi_h''}$$

durch p teilbar, und F_h erhält daher einen Wert

$$\equiv \Phi_h' \left(1 \pm \frac{\Phi_h'''}{\sigma_h \Phi_h'} \right)^2 \quad \text{oder} \quad \equiv \Phi_h'' \left(1 \pm \frac{\Phi_h'''}{\sigma_h \Phi_h''} \right)^2 \pmod{p}.$$

In beiden Fällen besitzt offenbar F_h keinen andern quadratischen Charakter in bezug auf p als die Zahlen Φ_h .

Sobald $\sigma_{h-1} o_h \sigma_{h+1} \equiv 4 \pmod{8}$ oder $\equiv 0 \pmod{8}$ wird, ist die Zahl σ_h immer gleich 1. Wird hier eine der Zahlen Φ_h' , Φ_h'' ungerade, so ist das betreffende F_h nach dem Modul 4 oder 8 entweder

$$\equiv \Phi_h' \left(1 \pm \frac{\Phi_h'''}{\sigma_h \Phi_h'} \right)^2 \quad \text{oder} \quad \equiv \Phi_h'' \left(1 \pm \frac{\Phi_h'''}{\sigma_h \Phi_h''} \right)^2;$$

wenn dagegen Φ_h' sowohl als Φ_h'' kongruent Null $\pmod{4}$ wird, so ist wegen $\Phi_h' \Phi_h'' - (\Phi_h''')^2 \equiv 0 \pmod{\sigma_{h-1} o_h \sigma_{h+1}}$ Φ_h''' gerade, und der Ausdruck $F_h = \Phi_h' \pm 2 \Phi_h''' + \Phi_h''$ erhält gleichfalls einen Wert $\equiv 0 \pmod{4}$. Demnach besitzen die F_h immer dieselben quadratischen Charaktere für den Modul 4 oder 8 wie die Φ_h .

Aus dem Satze II (Kap. V) ersehen wir noch, daß die Charaktere, welche wir soeben betrachtet haben, leicht aus einer jeden vorgelegten Form f erschlossen werden können, ohne daß es nötig wäre, f einer linearen Transformation zu unterwerfen. Denn sobald $o_h \equiv 0 \pmod{p}$ ist, gibt es unter den Größen $\Delta_h(f)$, welche ja selbst Zahlen (f_h) sind, solche, die zu p relativ prim werden und für die demnach $\left(\frac{\Delta_h(f)}{p} \right) = \left(\frac{f_h}{p} \right)$ ist; und ähnlich gibt es in den Fällen $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{4}$ oder $\equiv 0 \pmod{8}$ unter den Größen $\Delta_h(f)$ immer ungerade, für welche also resp.

$$(-1)^{\frac{\Delta_h(f)-1}{2}} = (-1)^{\frac{(f_h)-1}{2}} \quad \text{und} \quad \left(\frac{2}{\Delta_h(f)} \right) = \left(\frac{2}{(f_h)} \right)$$

wird.

Eine zweite Methode zur Herleitung der Charaktere

$$\left(\frac{f_h}{p}\right), \quad (-1)^{\frac{(f_h)-1}{2}}, \quad \left(\frac{2}{f_h}\right)$$

ergibt sich unmittelbar aus den Formeln (2) und (3) des Kap. V.

Außer den Charakteren, die wir bisher gefunden haben, gibt es noch eine Reihe weiterer Charaktere, deren Existenz *mit Benutzung des quadratischen Reziprozitätsgesetzes* durch ein analoges Verfahren nachgewiesen werden kann. Wir behalten uns vor, die hierauf bezüglichen Bemerkungen bei einer anderen Gelegenheit ausführlicher darzulegen.

Jetzt wenden wir uns dazu, den Fall $h = 1$ der Kongruenzen $(f_h) \equiv m_h \pmod{N_0}$ zu diskutieren.

Kap. VII. Über die Anzahl der Lösungen der Kongruenzen

$$f = \sum_{i,k=1}^n a_{ik} x_i x_k \equiv m \pmod{N}.$$

I. Wir bezeichnen die Anzahl sämtlicher für den Modul N inkongruenten Lösungssysteme (X_i) der Kongruenz

$$f = \sum_{i,k=1}^n a_{ik} X_i X_k \equiv m \pmod{N}$$

durch $f\{m; N\}$.

Gehören die Formen f und g zu derselben Gruppe nach einem Modul N , so wird

$$(9) \quad f\{m; N\} = g\{m; N\}.$$

Denn es möge die Form f durch die Substitution $x_i = \sum_{k=1}^n s_k^i y_k$ in eine äquivalente Form $f_0 \equiv g \pmod{N}$ übergehen. Wir gewinnen dann aus einem jeden System $(Y_k) \pmod{N}$, für welches $g(Y_k) \equiv f_0(Y_k) \equiv m \pmod{N}$ wird, ein bestimmtes System $X_i \equiv \sum_{k=1}^n s_k^i Y_k \pmod{N}$, für welches

$f(X_i) \equiv m \pmod{N}$ wird, und aus zwei modulo N verschiedenen Systemen (Y_k) gewinnen wir wegen $|s_k^i| = 1$ stets zwei modulo N verschiedene Systeme (X_i) . Folglich wird $f\{m; N\} \geq g\{m; N\}$ sein. Aber auf dieselbe Weise schließt man $f\{m; N\} \leq g\{m; N\}$, woraus sich die zu beweisende Gleichung ergibt.

Die verschiedenen Zahlen $f\{m; N\}$, welche einer bestimmten Form f angehören, hängen einestheils von der Ordnung, andernteils von den Charakteren dieser Form ab. Um dies zu zeigen, tun wir gut, anstelle der

Größen $f\{m; N\}$, welche im allgemeinen ziemlich kompliziert ausfallen, Verbindungen derselben einzuführen, welche sich in einer einfacheren und übersichtlicheren Weise ausdrücken lassen. Zu solchen Verbindungen gelangen wir durch Hinzuziehung von Einheitswurzeln.

Es möge die Größe

$$e^{\frac{2\pi i}{N}} = \cos \frac{2\pi}{N} + i \sin \frac{2\pi}{N},$$

welche eine primitive N^{te} Einheitswurzel vorstellt, gleich r_N gesetzt werden. Wir wollen die Summen

$$(10) \quad f\{1; N\} r_N^{1 \cdot h} + f\{2; N\} r_N^{2 \cdot h} + \cdots + f\{N; N\} r_N^{N \cdot h} \\ = \sum_{m=1}^N f\{m; N\} r_N^{m \cdot h} = f(h; N)$$

betrachten, in welchen h eine ganze Zahl bedeutet.

Infolge der Gleichung $r_N^N = 1$ nehmen diese Summen $f(h; N)$ für zwei nach dem Modul N kongruente Zahlen h gleiche Werte an. Erteilen wir der Größe h die N für den Modul N inkongruenten Werte $1, 2, \dots, N$, so entstehen im ganzen N Größen

$$f(1; N), f(2; N), \dots, f(N; N),$$

durch welche sich umgekehrt die Zahlen $f\{m; N\}$ ausdrücken lassen. Wir bekommen

$$(11) \quad f\{m; N\} = \frac{1}{N} \cdot \sum_{h=1}^N f(h; N) r_N^{-h \cdot m}.$$

Die Gleichungen (10) und (11) zeigen, daß die Untersuchung der Größen $f\{m; N\}$ vollständig auf eine Untersuchung der Größen $f(h; N)$ hinauskommt.

Für zwei nach dem Modul N kongruente Formen f und g liefert die Gleichung (9) die Relationen $f(h; N) = g(h; N)$. Beachten wir, daß eine jede Form g , welche der Kongruenz $g \simeq f \pmod{N}$ genügt, einen Rest der Klasse f in bezug auf den Modul N vorstellt, so können wir den Satz aussprechen:

H. Ist φ ein beliebiger Rest der Klasse f , so gelten die Gleichungen

$$f(h; N) = \varphi(h; N) \quad \text{und} \quad f\{m; N\} = \varphi\{m; N\}.$$

Es hat die Identität

$$(12) \quad \sum_{m=1}^N f\{m; N\} r_N^{m \cdot h} = \sum_{\xi_i}^{1, N} r_N^{h \cdot f(\xi_i)}$$

statt, sobald in der zweiten Summe jede der Variablen ξ_i ein vollständiges Restsystem nach dem Modul N durchläuft. — Denn erteilen wir einer

jeden Größe ξ_i die sämtlichen N Werte $\equiv 1, 2, \dots, N \pmod{N}$, so nimmt $f(\xi_i)$ je $f\{1; N\}$ -mal einen Wert $\equiv 1 \pmod{N}$, je $f\{2; N\}$ -mal einen Wert $\equiv 2 \pmod{N}$, usf., je $f\{N; N\}$ -mal einen Wert $\equiv N \pmod{N}$ an.

In der Summe $\sum_{\xi_i}^{1, N} r_N^{hf(\xi_i)}$ wird demnach je $f\{1; N\}$ -mal ein Glied $r_N^{1 \cdot h}$, je $f\{2; N\}$ -mal ein Glied $r_N^{2 \cdot h}$, usf., je $f\{N; N\}$ -mal ein Glied $r_N^{N \cdot h}$ auftreten, und daraus ergibt sich die Gleichung (12).

Die Identität (12) liefert, mit der Gleichung (10) verbunden, eine neue Definition der Größen $f(h; N)$, nämlich

$$f(h; N) = \sum_{\xi_i}^{1, N} r_N^{hf(\xi_i) \cdot *}$$

Aus dieser Gleichung fließt leicht der folgende Hauptsatz, durch welchen die Bestimmung dieser Größen außerordentlich erleichtert wird:

J. Wenn die quadratische Form f nach dem Modul N in eine Summe von quadratischen Formen f_s [[mit getrennten Variablen]] zerfällt:

$$f(\xi_i) \equiv \sum_{s=1}^{s_0} f_s(\xi_k^{(s)}) \pmod{N},$$

so wird

$$f(h; N) = \prod_{s=1}^{s_0} f_s(h; N).$$

Denn bezeichnen wir die Variablen der einzelnen Formen f_s mit $\xi_k^{(s)}$, so erhalten wir die Relationen

$$f_s(h; N) = \sum_{\xi_k^{(s)}}^{1, N} r_N^{hf_s(\xi_k^{(s)})} \quad (s = 1, 2, \dots, s_0),$$

deren Multiplikation unmittelbar

$$\sum_{\xi_k^{(s)}}^{1, N} r_N^{h \sum_{s=1}^{s_0} f_s(\xi_k^{(s)})} = \sum_{\xi_i}^{1, N} r_N^{hf(\xi_i)} = f(h; N)$$

ergibt.

II. Die Größen $f(h; N)$ sind infolge der Beziehung $r_N^N = 1$ unabhängig von den besonderen Restsystemen, welche von den einzelnen Variablen ξ_i durchlaufen werden. Wir schließen aus diesem Umstande die Sätze:

1) Wenn N_0 ein Teiler von N ist, so gilt

$$(13) \quad f(h; N_0) = \left(\frac{N}{N_0}\right)^{-n} \cdot f\left(h \frac{N}{N_0}; N\right).$$

*) Die Summen $f(h; N)$ sind von Herrn Weber im Band 74 des Crelleschen Journals behandelt worden.

Bekanntlich setzt sich ein jedes vollständige Restsystem für den Modul N aus $\frac{N}{N_0}$ vollständigen Restsystemen für den Modul N_0 zusammen. Es ist daher

$$\sum_{\xi_i}^{1, N} r_{\frac{N}{N_0}}^{h \frac{N}{N_0} f(\xi_i)} = \left(\frac{N}{N_0}\right)^n \cdot \sum_{\xi_i}^{1, N_0} r_{N_0}^{h f(\xi_i)},$$

und hieraus geht sofort die Gleichung (13) hervor. Wenn insbesondere N eine Primzahlpotenz q^t ist, so folgt für $t \geq t_0$

$$f(h; q^t) = q^{-n(t-t_0)} \cdot f(hq^{t-t_0}; q^t).$$

2) Es ist

$$f(hZ^2; N) = f(h; N),$$

sobald Z eine zu N relativ prime Zahl bedeutet.

Wegen $f(Zx_i) = Z^2 f(x_i)$ ist die Größe $f(hZ^2; N)$ gleich der Summe

$$\sum_{x_i}^{1, N} r_{\frac{N}{N_0}}^{h f(Zx_i)}.$$

Setzen wir nun $Zx_i \equiv \xi_i \pmod{N}$, so werden offenbar die Zahlen ξ_i zugleich mit den Zahlen x_i vollständige Restsysteme nach dem Modul N durchlaufen. Also wird diese Summe gleich $f(h; N)$ sein.

3) Wenn die Zahl N ein Produkt aus c_0 zueinander relativ primen Zahlen N_c ist, so wird

$$f(h; N) = \prod_{c=1}^{c_0} f\left(h \frac{N}{N_c}; N_c\right).$$

Beweis: Wir wollen die Zahlen $\frac{N}{N_c} = L_c$ setzen. Jede dieser Zahlen L_c ist offenbar zu der entsprechenden Zahl N_c relativ prim, während sie durch alle übrigen Größen N_{c^*} ($c^* \neq c$) teilbar ist. Wenn wir in dem Ausdrucke

$$\xi \equiv L_1 \xi^{(1)} + L_2 \xi^{(2)} + \dots + L_{c_0} \xi^{(c_0)} \pmod{N}$$

die Größen $\xi^{(c)}$ vollständige Restsysteme nach den Moduln N_c durchlaufen lassen, so erhalten wir N Zahlen, welche ein vollständiges Restsystem für den Modul N bilden. Denn es können nicht zwei von diesen N Zahlen nach dem Modul N kongruent sein, da aus einer jeden Kongruenz

$$\sum L_c \xi^{(c)} \equiv \sum L_c \xi_0^{(c)} \pmod{N}$$

auf der Stelle

$$L_c \xi^{(c)} \equiv L_c \xi_0^{(c)} \pmod{N_c}$$

oder, weil die Zahlen L_c zu den Größen N_c relativ prim sind,

$$\xi^{(c)} \equiv \xi_0^{(c)} \pmod{N_c}$$

folgen würde. Wir dürfen demnach für die Größe $f(h; N)$ schreiben

$$f(h; N) = \sum_{\xi_i^{(c)}}^{1, N_c} r_{\frac{N}{N_c}}^{h \cdot f\left(\sum_{c=1}^{c_0} L_c \xi_i^{(c)}\right)}.$$

Durch Multiplikation der beiden Kongruenzen

$$\xi_i \equiv \sum_{c=1}^{c_0} L_c \xi_i^{(c)} \pmod{N} \quad \text{und} \quad \xi_k \equiv \sum_{c=1}^{c_0} L_c \xi_k^{(c)} \pmod{N}$$

erhalten wir, da das Produkt $L_c L_{c^*}$ für zwei ungleiche Indizes c und c^* gleich $N \cdot \frac{N}{N_c N_{c^*}} \equiv 0 \pmod{N}$ wird,

$$\xi_i \xi_k \equiv \sum_{c=1}^{c_0} L_c^2 \xi_i^{(c)} \xi_k^{(c)} \pmod{N},$$

also

$$f\left(\sum_{c=1}^{c_0} L_c \xi_i^{(c)}\right) \equiv \sum_{c=1}^{c_0} L_c^2 f(\xi_i^{(c)}) \pmod{N}.$$

Demnach bekommen wir für $f(h; N)$ die Gleichung

$$f(h; N) = \sum_{\xi_i^{(c)}}^{1, N_c} r_N^h \sum_{c=1}^{c_0} L_c^2 f(\xi_i^{(c)}) = \prod_{c=1}^{c_0} \left(\sum_{\xi_i^{(c)}}^{1, N_c} r_N^{L_c^2} f(\xi_i^{(c)}) \right).$$

Die Einzelglieder des Produkts in dieser Formel sind jetzt wegen $r_N^{L_c} = r_{N_c}$ gleich $f(h L_c; N_c)$, und wir gewinnen wirklich die Identität, welche wir beweisen wollten.

Wählen wir für die Zahlen N_c die verschiedenen in N enthaltenen Primzahlpotenzen q^t , so können mit Hilfe des soeben bewiesenen Satzes die Größen $f(h; N)$ auf Größen $f(h; q^t)$ zurückgeführt werden. Es erhellt hieraus sofort, daß wir uns weiterhin auf die Betrachtung von Größen $f(h; q^t)$ für Primzahlpotenzmoduln q^t beschränken dürfen. Wir brauchen ferner nur solche Moduln zu untersuchen, welche gewisse Grenzen $q^{G(q)}$ überschreiten. Denn es lassen sich mit Hilfe des Satzes 1) die Größen $f(h; q^t)$ für Moduln $q^t < q^t$ stets durch Größen $f(h; q^t)$ ausdrücken.

III. Um die Größen $f(h; q^t)$ für einen Primzahlmodul q^t zu finden, gehen wir von einem Hauptrest φ der Klasse f für den Modul q^t aus. Die diesem Reste φ angehörigen Größen $\varphi(h; q^t)$ sind nach dem Satz H. den Größen $f(h; q^t)$ gleich.

($q = p$). Wenn nun zunächst $q^t = p^t$ ist, so zerfällt der Hauptrest $\varphi \pmod{p^t}$ in eine Summe von Einzelformen von der Art

$$F \equiv p^M \alpha \xi^2 \pmod{p^t}, \quad [\alpha \text{ relativ prim zu } p]$$

und die Größen $\varphi(h; p^t)$ können mit Hilfe des Satzes J. durch die Größen

$$F(h; p^t) = \sum_{\xi=1}^{p^t} e^{\frac{2\pi i p^M h \alpha \xi^2}{p^t}} = \sum_{\xi=1}^{p^t} e^{\frac{2\pi i \tau \xi^2}{p^t}} = (\tau; p^t) \quad [\tau = p^M h \alpha]$$

dargestellt werden.

($q = 2$). Wenn hingegen $q^t = 2^t$ ist, so zerfällt der Hauptrest $\varphi \pmod{2^t}$ in eine Summe von Einzelformen von der Art

$$F \equiv 2^M \alpha \xi^2 \pmod{2^t} \quad [\alpha \text{ relativ prim zu } 2]$$

oder von der Art

$$F_0 \equiv 2^M (\alpha \xi^2 + \mathfrak{A} \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2) \pmod{2^t} \quad [\mathfrak{A} \equiv 1 \pmod{2}],$$

und die Größen $\varphi(h; 2^t)$ lassen sich mit Hilfe des Satzes J. durch die Größen

$$F(h; 2^t) = \sum_{\xi=1}^{2^t} e^{\frac{2\pi i 2^M h \alpha \xi^2}{2^t}} = \sum_{\xi=1}^{2^t} e^{\frac{2\pi i \tau \xi^2}{2^t}} = (\tau; 2^t) \quad [\tau = 2^M h \alpha]$$

und durch Größen

$$F_0(h; 2^t) = \sum_{\xi, \tilde{\xi}=1}^{2^t} e^{\frac{2\pi i h 2^M}{2^t} (\alpha \xi^2 + \mathfrak{A} \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2)}$$

darstellen.

Kap. VIII. Bestimmung der Größen $f(h; q^t)$ in den einfachsten Fällen.

I. Vermittels der soeben aufgestellten Sätze können die allgemeinen Größen $f(h; q^t)$ stets auf die drei Summen

$$(\tau; p^t), \quad (\tau; 2^t), \quad F_0(h; 2^t)$$

zurückgeführt werden. Wir wollen jetzt der Reihe nach die Werte dieser Summen aufsuchen.

($q = p$). Wenn τ relativ prim zu p ist, so beweist man leicht die Beziehung

$$(\tau; p) = \left(\frac{\tau}{p}\right) (1; p).$$

Um $(1; p)$ zu erhalten, betrachten wir einen Rest

$$\Phi \equiv x_1 x_2 \equiv \begin{pmatrix} 0, & 1 \\ 1, & 0 \end{pmatrix} \pmod{p},$$

für welchen

$$\Phi(1; p) = \sum_{x_1=1}^p \sum_{x_2=1}^p e^{\frac{2\pi i x_1 x_2}{p}} = p$$

ist. Die Form Φ besitzt nach unseren Sätzen einen Hauptrepräsentanten

$$\Phi_0 \equiv \begin{pmatrix} a, & 0 \\ 0, & a_0 \end{pmatrix} \equiv ax^2 + a_0 x_0^2 \pmod{p},$$

in welchem a und a_0 zu p relativ prim sind. Die Determinante dieses Repräsentanten Φ_0 wird mit der Determinante von Φ nach dem Modul p übereinstimmen müssen; wir erhalten daher

$$aa_0 \equiv -1 \pmod{p},$$

und es ist für die Form Φ_0

$$\Phi_0(1; p) = \sum_{x=1}^p e^{\frac{2\pi i a x^2}{p}} \cdot \sum_{x_0=1}^p e^{\frac{2\pi i a_0 x_0^2}{p}} = \left(\frac{aa_0}{p}\right) (1; p)^2 = (-1)^{\frac{p-1}{2}} (1; p)^2.$$

Da nun wegen $\Phi \equiv \Phi_0 \pmod{p}$ die Größen $\Phi(1; p)$ und $\Phi_0(1; p)$ identisch sein müssen, so bekommen wir die Gleichung

$$(1; p)^2 = (-1)^{\frac{p-1}{2}} p,$$

und wir können mithin

$$(1; p) = i^{\left(\frac{p-1}{2}\right)^2} p^{\frac{1}{2}} \delta_p$$

setzen, wo die Größe δ_p eine nur von p abhängige Einheit bezeichnet. Bekanntlich ist diese Einheit stets gleich $+1$.

Vermittels der Größe $(1; p)$ kann man leicht die allgemeine Summe $(\tau; p^t)$ ausdrücken, und wenn man $\tau = p^T \tau_0$ (τ_0 relativ prim zu p) setzt, erhält man die Gleichungen

$$(14_0) \quad \text{für } t - T \leq 0: \quad (\tau; p^t) = p^t,$$

$$(14_1) \quad \text{für } t - T > 0: \quad (\tau; p^t) = \left(\frac{\tau_0}{p}\right)^{t-T} p^{\frac{t+T}{2}} i^{\frac{1}{4}(p^{t-T}-1)^2} \delta_p^{t-T}.$$

($q=2$). Für die Summe $(\tau; 2^t)$ gelten, wenn man $\tau = 2^T \tau_0 [\tau_0 \equiv 1 \pmod{2}]$ setzt, die Formeln

$$(15_0) \quad \text{für } t - T < 1: \quad (\tau; 2^t) = 2^t,$$

$$(15_1) \quad \text{für } t - T = 1: \quad (\tau; 2^t) = 0,$$

$$(15_2) \quad \text{für } t - T > 1: \quad (\tau; 2^t) = \left(1 + i(-1)^{\frac{\tau_0-1}{2}}\right) \left(\frac{2}{\tau_0}\right)^{t-T} 2^{\frac{t+T}{2}}.*$$

Mit Hilfe der Summen $(\tau; 2^t)$ können wir jetzt den Wert der Summe

$$\Sigma = \sum_{\xi, \tilde{\xi}=1}^{2^t} e^{\frac{2\pi i \tau}{2^t} (\alpha \xi^2 + \mathfrak{A} \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2)} \quad [\mathfrak{A} \equiv 1 \pmod{2}]$$

bestimmen. Wir wollen die Größe $4\alpha\tilde{\alpha} - \mathfrak{A}^2$ durch D bezeichnen.

Es möge zunächst τ ungerade und $t \geq 1$ sein. Für einen Rest von der Art

$$\Phi \equiv \xi_0^2 + 2\tau(\alpha \xi^2 + \mathfrak{A} \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2) \equiv \begin{pmatrix} 1, & 0, & 0 \\ 0, & 2\tau\alpha, & \tau\mathfrak{A} \\ 0, & \tau\mathfrak{A}, & 2\tau\tilde{\alpha} \end{pmatrix} \pmod{2^{t+1}}$$

*) Gauß, Summatio quarundam serierum singularium. Gesammelte Werke, Bd. II.

wird

$$\Phi(1; 2^{t+1}) = 4 \sum_{\xi_0=1}^{2^{t+1}} e^{\frac{2\pi i \xi_0^2}{2^{t+1}}} \cdot \sum_{\xi, \bar{\xi}=1}^{2^t} e^{\frac{2\pi i \tau}{2^t} (\alpha \xi^2 + \mathfrak{A} \xi \bar{\xi} + \bar{\alpha} \bar{\xi}^2)} = 4 \cdot (1; 2^{t+1}) \Sigma.$$

Die Summe $(1; 2^{t+1})$ erhält nach der Gleichung (15₂), da $t+1$ nach unserer Annahme größer als 1 ist, den Wert $(1+i)2^{\frac{t+1}{2}}$, und wir bekommen daher

$$\Phi(1; 2^{t+1}) = 2^{\frac{t+5}{2}} (1+i) \Sigma.$$

Die Form Φ besitzt, da ihre Determinante ungerade ist und ihre beiden Invarianten σ gleich 1 sind, für den Modul 2^{t+1} einen Hauptrepräsentanten von der Art

$$\Phi_0 \equiv \begin{pmatrix} a_1, 0, 0 \\ 0, a_2, 0 \\ 0, 0, a_3 \end{pmatrix} \equiv a_1 x_1^2 + a_2 x_2^2 + a_3 x_3^2 \pmod{2^{t+1}},$$

in welchem die Größen a_1, a_2, a_3 ungerade sind. Da die Determinante der Form Φ_0 der Determinante von Φ für den Modul 2^{t+1} kongruent wird, so besteht die Kongruenz

$$a_1 a_2 a_3 \equiv D \cdot \tau^2 \pmod{2^{t+1}},$$

aus welcher

$$a_1 a_2 a_3 \equiv -1 \pmod{4}$$

folgt.

Es müssen demnach von den Größen a_1, a_2, a_3 entweder eine oder drei $\equiv -1 \pmod{4}$ sein; wir erkennen leicht, daß der erste oder der zweite Fall eintritt, je nachdem die Einheit

$$\delta = (-1)^{\frac{a_2-1}{2} \cdot \frac{a_3-1}{2} + \frac{a_3-1}{2} \cdot \frac{a_1-1}{2} + \frac{a_1-1}{2} \cdot \frac{a_2-1}{2}}$$

gleich $+1$ oder gleich -1 ist. Um zu entscheiden, welchen Wert diese Einheit annimmt, beachten wir, daß die Kongruenzen $\Phi \equiv 1 \pmod{4}$ und $\Phi_0 \equiv 1 \pmod{4}$ wegen $\Phi \subseteq \Phi_0 \pmod{2^{t+1} \geq 4}$ gleich viel Lösungen zulassen müssen. Nun finden wir die Anzahl der Lösungen der Kongruenz $\Phi \equiv 1 \pmod{4}$ gleich $2^4 \left[1 + \frac{1}{2} \left(\frac{2}{D} \right) \right]$ und die Anzahl der Lösungen von $\Phi_0 \equiv 1 \pmod{4}$ gleich $2^4 \left(1 + \frac{1}{2} \delta \right)$; es wird daher die Beziehung

$$\delta = \left(\frac{2}{D} \right)$$

bestehen. Für den Formenrest Φ_0 ergibt sich

$$\Phi_0(1; 2^{t+1}) = (a_1; 2^{t+1}) \cdot (a_2; 2^{t+1}) \cdot (a_3; 2^{t+1}).$$

Mit Hilfe der Formel (15₂) schließen wir hieraus, da $t+1 > 1$ ist,

$$\Phi_0(1; 2^{t+1}) = \left[1 + i(-1)^{\frac{a_1-1}{2}}\right] \left[1 + i(-1)^{\frac{a_2-1}{2}}\right] \left[1 + i(-1)^{\frac{a_3-1}{2}}\right] \cdot \left(\frac{2}{a_1 a_2 a_3}\right)^{t+1} 2^{\frac{3t+3}{2}}.$$

Aus der Kongruenz $D \cdot \tau^2 \equiv a_1 a_2 a_3 \pmod{2^{t+1} \geq 4}$ gewinnen wir die Gleichung

$$\left(\frac{2}{a_1 a_2 a_3}\right)^{t+1} = \left(\frac{2}{D}\right)^{t+1};$$

ferner finden wir

$$\left[1 + i(-1)^{\frac{a_1-1}{2}}\right] \left[1 + i(-1)^{\frac{a_2-1}{2}}\right] \left[1 + i(-1)^{\frac{a_3-1}{2}}\right] = 2(1+i)\delta = 2(1+i)\left(\frac{2}{D}\right)$$

Also wird

$$\Phi_0(1; 2^{t+1}) = \left(\frac{2}{D}\right)^t (1+i) 2^{\frac{3t+5}{2}}.$$

Da nun wegen $\Phi \equiv \Phi_0 \pmod{2^{t+1}}$ die Größen $\Phi(1; 2^{t+1})$ und $\Phi_0(1; 2^{t+1})$ einander gleich sind, so erhalten wir

$$\Sigma = \left(\frac{2}{D}\right)^t 2^t.$$

Wenn die Größe τ durch eine Potenz von 2 teilbar ist, so können wir τ in der Form $\tau = 2^T \tau_0 [\tau_0 \equiv 1 \pmod{2}]$ darstellen, und wir bekommen alsdann die beiden Formeln

$$(16_0) \quad \text{für } t - T \leq 0: \quad \Sigma = 2^{2t},$$

$$(16_1) \quad \text{für } t - T > 0: \quad \Sigma = \left(\frac{2}{D}\right)^{t-T} 2^{t+T}.$$

II. Wir benutzen jetzt die gefundenen Summen zur Bestimmung von Größen $\varphi(h; q^t)$ für einige besonders einfache Reste $\varphi \pmod{q^t}$. Wir stellen die Zahl h in der Form $h = q^{\bar{h}} h_0$ (h_0 relativ prim zu q) dar.

($q = p$). Die Größen $F(h; p^t)$ eines Restes $F \equiv p^M \alpha \xi^2 \pmod{p^t}$ [α relativ prim zu p] sind durch die Gleichungen

$$F(h; p^t) = (h p^M \alpha; p^t)$$

gegeben. Es ergeben sich die Formeln

$$\text{für } \bar{h} \geq t - M: \quad F(h; p^t) = p^t,$$

$$\text{für } \bar{h} < t - M: \quad F(h; p^t) = \left(\frac{h_0 \alpha}{p}\right)^{t-M-\bar{h}} p^{\frac{t+M+\bar{h}}{2}} i^{\frac{1}{4}(p^{t-M-\bar{h}}-1)^2} \delta_p^{t-M-\bar{h}}.$$

Die Größen $\varphi(h; p^t)$, welche einem Rest

$$\varphi \equiv p^M (\alpha_1 \xi_1^2 + \alpha_2 \xi_2^2 + \cdots + \alpha_x \xi_x^2) \equiv \sum_{i=1}^x p^M \alpha_i \xi_i^2 \equiv \sum_{i=1}^x F_i \pmod{p^t}$$

angehören, genügen den Relationen

$$\varphi(h; p^t) = \prod_{i=1}^x F_i(h; p^t).$$

Setzen wir das Produkt $\prod_{i=1}^{\kappa} \alpha_i \equiv (\varphi) \pmod{p^{t-M}}$, so gewinnen wir die Formeln

$$\text{für } \bar{h} \geq t - M: \quad \varphi(h; p^t) = p^{\kappa t},$$

$$\text{für } \bar{h} < t - M:$$

$$\varphi(h; p^t) = \left(\frac{(\varphi)}{p}\right)^{t-M-\bar{h}} p^{\frac{\kappa(t+M+\bar{h})}{2}} \left(\frac{h_0}{p}\right)^{\kappa(t-M-\bar{h})} \frac{\kappa}{2^{\frac{1}{4}}} (p^{t-M-\bar{h}} - 1)^{\frac{\kappa}{2}} \delta_p^{\kappa(t-M-\bar{h})}.$$

Aus denselben ersehen wir, daß die Größen $\varphi(h; p^t)$ sich in zwei Faktoren zerlegen lassen, von denen der erste eine Einheit vorstellt, während der zweite dem absoluten Werte nach mit $\varphi(h; p^t)$ übereinstimmt und allein von den Größen $h, p^t; \kappa, M$ abhängt. Bezeichnen wir diesen zweiten Faktor durch $\overset{\circ}{\varphi}(h; p^t)$, so erhalten wir die Formeln

$$\text{für } \bar{h} \geq t - M: \quad \varphi(h; p^t) = \overset{\circ}{\varphi}(h; p^t),$$

$$\text{für } \bar{h} < t - M: \quad \varphi(h; p^t) = \left(\frac{(\varphi)}{p}\right)^{t-M-\bar{h}} \cdot \overset{\circ}{\varphi}(h; p^t).$$

($q = 2$). Die Größen $F(h; 2^t)$ eines Restes $F \equiv 2^M \alpha \xi^2 \pmod{2^t}$ [α relativ prim zu 2] haben die Werte

$$F(h; 2^t) = (h 2^M \alpha; 2^t),$$

und es ergeben sich die drei Formeln

$$\text{für } \bar{h} > t - 1 - M: \quad F(h; 2^t) = 2^t,$$

$$\text{für } \bar{h} = t - 1 - M: \quad F(h; 2^t) = 0,$$

$$\text{für } \bar{h} < t - 1 - M: \quad F(h; 2^t) = \left[1 + i(-1)^{\frac{h_0 \alpha - 1}{2}}\right] \left(\frac{2}{h_0 \alpha}\right)^{t-M-\bar{h}} 2^{\frac{t+M+\bar{h}}{2}}.$$

Die Größen $\varphi(h; 2^t)$ eines Restes

$$\varphi \equiv 2^M (\alpha_1 \xi_1^2 + \alpha_2 \xi_2^2 + \dots + \alpha_{\kappa} \xi_{\kappa}^2) = \sum_{i=1}^{\kappa} 2^M \alpha_i \xi_i^2 \equiv \sum_{i=1}^{\kappa} F_i \pmod{2^t}$$

genügen den Gleichungen

$$\varphi(h; 2^t) = \prod_{i=1}^{\kappa} F_i(h; 2^t).$$

Setzen wir das Produkt $\prod_{i=1}^{\kappa} \alpha_i \equiv (\varphi) \pmod{2^{t-M}}$, so folgen die Relationen

$$(17) \begin{cases} \text{für } \bar{h} > t - 1 - M: & \varphi(h; 2^t) = 2^{\kappa t}, \\ \text{für } \bar{h} = t - 1 - M: & \varphi(h; 2^t) = 0, \\ \text{für } \bar{h} < t - 1 - M: & \varphi(h; 2^t) = \prod_{i=1}^{\kappa} \left[1 + i(-1)^{\frac{h_0 \alpha_i - 1}{2}}\right] \cdot \left(\frac{2}{(\varphi)}\right)^{t-M-\bar{h}} \left(\frac{2}{h_0}\right)^{\kappa(t-M-\bar{h})} 2^{\frac{\kappa(t+M+\bar{h})}{2}}. \end{cases}$$

Die Größen $\Phi(h; 2^t)$ eines Restes

$$\begin{aligned}\Phi &\equiv 2^M \left(\sum_1^{l_1} \alpha_i^{(1)} \xi_i^{(1)} \xi_i^{(1)} + 2 \sum_1^{l_2} \alpha_i^{(2)} \xi_i^{(2)} \xi_i^{(2)} + \dots + 2^{m-1} \sum_1^{l_m} \alpha_i^{(m)} \xi_i^{(m)} \xi_i^{(m)} \right) \\ &\equiv \sum_{s=1}^m \left(2^{M+s-1} \sum_{i=1}^{l_s} \alpha_i^{(s)} \xi_i^{(s)} \xi_i^{(s)} \right) \equiv \sum_{s=1}^m \varphi_s \pmod{2^t}\end{aligned}$$

$$[l_s > 0, l_1 + l_2 + \dots + l_m = l]$$

sind gleich

$$\Phi(h; 2^t) = \prod_{s=1}^m \varphi_s(h; 2^t).$$

Aus den Beziehungen (17) schließen wir für die einzelnen Reste φ_s die Gleichungen

$$(18_0) \quad \text{für } \bar{h} > t - s - M: \quad \varphi_s(h; 2^t) = 2^{l_s t},$$

$$(18_1) \quad \text{für } \bar{h} = t - s - M: \quad \varphi_s(h; 2^t) = 0,$$

$$(18_2) \quad \text{für } \bar{h} < t - s - M:$$

$$\varphi_s(h; 2^t) = \prod_{i=1}^{l_s} \left[1 + i(-1)^{\frac{h_0 \alpha_i^{(s)} - 1}{2}} \right] \cdot \left(\frac{2}{(\varphi_s)} \right)^{t-M-s-1-\bar{h}} 2^{\frac{t+M+s-1+\bar{h}}{2}} \left(\frac{2}{h_0} \right)^{l_s(t-M-s-1-\bar{h})}.$$

Wenn $\bar{h} \geq t - M$ ist, so haben für die sämtlichen Größen $\varphi_s(h; 2^t)$ die Relationen (18₀) statt. In diesem Falle findet sich daher

$$(19_0) \quad \text{für } \bar{h} \geq t - M: \quad \Phi(h; 2^t) = 2^{lt}.$$

Befriedigt die Zahl $\bar{h}$ die Ungleichung $t - M > \bar{h} \geq t - m - M$, so können wir eine Zahl m_0 aus der Reihe $1, 2, \dots, m$ angeben, für welche $\bar{h} = t - m_0 - M$ ist. Dann tritt für die Größe $\varphi_{m_0}(h; 2^t)$ des Restes φ_{m_0} , welcher dieser Zahl m_0 entspricht, die Gleichung (18₁) in Kraft, und es verschwindet daher diese Größe $\varphi_{m_0}(h; 2^t)$. Demnach wird auch $\Phi(h; 2^t)$ gleich Null, und wir bekommen

$$(19_1) \quad \text{für } t - M > \bar{h} \geq t - m - M: \quad \Phi(h; 2^t) = 0.$$

Fällt die Zahl $\bar{h}$ kleiner als $t - m - M$ aus, so gelten für die sämtlichen Größen $\varphi_s(h; 2^t)$ die Gleichungen (18₂); es ergibt sich, wenn wir

$$\prod_{s=1}^m (\varphi_s) \equiv (\Phi) \pmod{2^{t-M-m+1}} \text{ setzen,}$$

$$(19_2) \quad \text{für } t - m - M > \bar{h}:$$

$$\begin{aligned}\Phi(h; 2^t) &= \prod_{s=1}^m \prod_{i=1}^{l_s} \left[1 + i(-1)^{\frac{h_0 \alpha_i^{(s)} - 1}{2}} \right] \cdot \left\{ \left(\frac{2}{(\varphi_{m-1})} \right) \left(\frac{2}{(\varphi_{m-2})} \right) \dots \left(\frac{2}{(\varphi_{m+1-2 \lfloor \frac{m}{2} \rfloor})} \right) \right\} \\ &\quad \cdot \left(\frac{2}{(\Phi)} \right)^{t-M-m-1-\bar{h}} 2^{\sum_{s=1}^m l_s \frac{t+M+s-1+\bar{h}}{2}} \left(\frac{2}{h_0} \right)^{\sum_{s=1}^m l_s (t-M-s-1-\bar{h})}\end{aligned}$$

Das in dieser letzten Formel auftretende Produkt

$$\prod \left[1 + i(-1)^{\frac{h_0 \alpha_i^{(s)} - 1}{2}} \right]$$

besitzt die Gestalt

$$\Pi = \prod_{k=1}^l \left[1 + i(-1)^{\frac{\delta_k - 1}{2}} \right].$$

Wir wollen jetzt nachsehen, von welchen Argumenten ein derartiges Produkt eigentlich abhängt. — Nehmen wir an, von den Größen $(-1)^{\frac{\delta_k - 1}{2}}$ seien im ganzen l_0 gleich -1 und $l - l_0$ gleich $+1$, so erhält Π den Wert

$$\Pi = (1 + i)^{l - l_0} (1 - i)^{l_0}$$

oder wegen der Beziehung $(1 - i) = -i(1 + i)$ den Wert

$$\Pi = (-1)^{l_0} i^{l_0} (1 + i)^l.$$

Nun ist

$$i^{l_0} = i^{2 \left[\frac{l_0}{2} \right]} i^{l_0 - 2 \left[\frac{l_0}{2} \right]} = (-1)^{\left[\frac{l_0}{2} \right]} i^{l_0 - 2 \left[\frac{l_0}{2} \right]},$$

und die Größe $i^{l_0 - 2 \left[\frac{l_0}{2} \right]}$ wird offenbar gleich 1 oder gleich i , je nachdem l_0 gerade oder ungerade ist. Wir können daher schreiben

$$i^{l_0 - 2 \left[\frac{l_0}{2} \right]} = \frac{1 + (-1)^{l_0}}{2} + \frac{1 - (-1)^{l_0}}{2} i = i^{l_0^2},$$

und wir bekommen die Gleichung

$$\Pi = (-1)^{\left[\frac{l_0}{2} \right]} \left[\frac{1 + (-1)^{l_0}}{2} - \frac{1 - (-1)^{l_0}}{2} i \right] (1 + i)^l.$$

Dieselbe zeigt uns, daß einerseits das Produkt Π durch die Zahl l und die beiden Einheiten $(-1)^{l_0}$ und $(-1)^{\left[\frac{l_0}{2} \right]}$ vollständig bestimmt ist, während andererseits aus dem Werte von Π die Werte der Größen l , $(-1)^{l_0}$, $(-1)^{\left[\frac{l_0}{2} \right]}$ erschlossen werden können.

Die Einheiten $(-1)^{l_0}$, $(-1)^{\left[\frac{l_0}{2} \right]}$ lassen sich in der folgenden Weise durch die Einheiten $(-1)^{\frac{\delta_k - 1}{2}}$ ausdrücken:

$$(-1)^{l_0} = (-1)^{\sum_{k=1}^l \frac{\delta_k - 1}{2}}, \quad (-1)^{\left[\frac{l_0}{2} \right]} = (-1)^{\sum_{k < k'}^{l, l} \frac{\delta_{k'} - 1}{2} \cdot \frac{\delta_{k''} - 1}{2}}.$$

Denn da erstens $\frac{\delta_k - 1}{2}$ gerade oder ungerade wird, je nachdem $(-1)^{\frac{\delta_k - 1}{2}}$ gleich $+1$ oder gleich -1 ist, so erscheinen in der Summe $\sum \frac{\delta_k - 1}{2}$,

welche sich über alle Zahlen $\frac{\delta_k - 1}{2}$ erstreckt, im ganzen $l - l_0$ gerade und l_0 ungerade Glieder; die Summe ist daher $\equiv l_0 \pmod{2}$. — Zweitens wird die Größe $\frac{\delta_{k'} - 1}{2} \cdot \frac{\delta_{k''} - 1}{2}$ ungerade oder gerade, je nachdem die Zahlen $(-1)^{\frac{\delta_{k'} - 1}{2}}$ und $(-1)^{\frac{\delta_{k''} - 1}{2}}$ alle beide gleich -1 sind oder das Gegenteil stattfindet. Bilden wir daher die Summe $\sum \frac{\delta_{k'} - 1}{2} \cdot \frac{\delta_{k''} - 1}{2}$ über alle möglichen Kombinationen zweier verschiedener Indizes k' und k'' , so treten in dieser Summe genau soviel ungerade Glieder auf, als Verbindungen von je zwei Größen $(-1)^{\frac{\delta_{k'} - 1}{2}} = -1$ und $(-1)^{\frac{\delta_{k''} - 1}{2}} = -1$ existieren. Unter den l_0 Zahlen $(-1)^{\frac{\delta_{k'} - 1}{2}}$, welche gleich -1 sind, gibt es nun offenbar im ganzen $l_0 \frac{l_0 - 1}{2} \equiv \left[\frac{l_0}{2} \right] \pmod{2}$ derartige Verbindungen, und es wird demnach

$$\sum_{k' < k''}^{1, l} \frac{\delta_{k'} - 1}{2} \cdot \frac{\delta_{k''} - 1}{2} \equiv \left[\frac{l_0}{2} \right] \pmod{2}.$$

Anstatt der Formel (19₂) erhalten wir jetzt für $t - m - M > \bar{h}$:

$$(19_3) \left\{ \begin{aligned} \Phi(h; 2^t) &= \left[\frac{1 + (-1)^{\frac{(\Phi) - 1}{2}}}{2} - \frac{1 - (-1)^{\frac{(\Phi) - 1}{2}}}{2} i \right] (-1)^{\frac{(\Phi) - 1}{2} \cdot \frac{h_0 - 1}{2}} \\ &\cdot \left(\frac{2}{(\Phi)} \right)^{t - M - m - 1 - \bar{h}} \left(\frac{2}{(\varphi_{m-1})} \right) \left(\frac{2}{(\varphi_{m-3})} \right) \cdots \left(\frac{2}{\left(\varphi_{m+1-2 \left[\frac{m}{2} \right]} \right)} \right) \\ &\cdot (-1)^{\frac{1}{2} \cdot \sum_{(s', s'') \neq (s'', s')} \frac{\alpha_{s'} - 1}{2} \cdot \frac{\alpha_{s''} - 1}{2}} (1 + i)^l 2^{\sum_{s=1}^m l_s \frac{t + M + s - 1 + \bar{h}}{2}} (-i)^l \left(\frac{h_0 - 1}{2} \right)^2 \\ &\cdot \sum_{s=1}^m l_s (t - M - s - 1 - \bar{h}) \\ &\cdot \left(\frac{2}{h_0} \right) \end{aligned} \right.$$

Wir sehen, daß alle nicht-verschwindenden Größen $\Phi(h; 2^t)$ sich in zwei Faktoren zerlegen lassen, von denen der erste eine Potenz von $i = \sqrt{-1}$ ist, während der zweite dem absoluten Werte nach mit $\Phi(h; 2^t)$ übereinstimmt und allein von den Größen $h, 2^t, m, l_s, M$ abhängt. Setzen wir diesen zweiten Faktor gleich $\hat{\Phi}(h; 2^t)$, so gewinnen wir die Formeln

$$(20_0) \quad \text{für } \bar{h} \geq t - M: \quad \Phi(h; 2^t) = \hat{\Phi}(h; 2^t),$$

$$(20_1) \quad \text{für } t - M > \bar{h} \geq t - m - M: \quad \Phi(h; 2^t) = 0,$$

$$(20_2) \quad \text{für } t - m - M > \bar{h}:$$

$$\begin{aligned} \Phi(h; 2^t) &= \overset{0}{\Phi}(h; 2^t) \left[\frac{1 + \frac{(-1)^{\frac{(\Phi)-1}{2}}}{2}}{2} - \frac{1 - \frac{(-1)^{\frac{(\Phi)-1}{2}}}{2}}{2} i \right] \cdot \\ &\cdot (-1)^{\frac{(\Phi)-1}{2} \cdot \frac{h_0-1}{2}} \left(\frac{2}{(\Phi)} \right)^{t-M-m-1-\bar{h}} \left(\frac{2}{(\varphi_{m-1})} \right) \left(\frac{2}{(\varphi_{m-3})} \right) \cdots \left(\frac{2}{(\varphi_{m+1-2\left[\frac{m}{2}\right]})} \right) \\ &\cdot (-1)^{\frac{1}{2} \cdot \sum_{(s', s') \neq (s'', s'')} \frac{\alpha_{s'}^{s'} - 1}{2} \cdot \frac{\alpha_{s''}^{s''} - 1}{2}}. \end{aligned}$$

Von Bedeutung für die Folge ist die Bemerkung, daß sowohl im Falle $l=1$ als im Falle $l=2$, $(-1)^{\frac{(\Phi)-1}{2}} = -1$ die Einheit

$$(-1)^{\frac{1}{2} \sum \frac{\alpha_{s'}^{s'} - 1}{2} \cdot \frac{\alpha_{s''}^{s''} - 1}{2}}$$

gleich $+1$ gesetzt werden muß.

Für die Größen $\Phi(h; 2^t)$ eines Restes

$$\Phi \equiv 2^M(\alpha \xi^2 + \mathfrak{A} \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2) \pmod{2^t}$$

gelten, wenn wir $4\alpha\tilde{\alpha} - \mathfrak{A}^2 \equiv (\Phi) \pmod{2^{2+t-M}}$ setzen, die Formeln

$$\text{für } \bar{h} \geq t - M: \quad \Phi(h; 2^t) = 2^{2t} = \overset{0}{\Phi}(h; 2^t),$$

$$\text{für } \bar{h} < t - M:$$

$$\Phi(h; 2^t) = \left(\frac{2}{(\Phi)} \right)^{t-M-\bar{h}} 2^{t+M+\bar{h}} = \overset{0}{\Phi}(h; 2^t) \cdot \left(\frac{2}{(\Phi)} \right)^{t-M-\bar{h}},$$

worin die Größen $\overset{0}{\Phi}(h; 2^t)$ allein von den Zahlen $h, 2^t; M$ abhängen.

Kap. IX. Charaktere der Hauptrepräsentanten und der Grundformen.

Wir werden jetzt nachweisen, daß die Größen $f(h; q^t)$ einer allgemeinen Form f ebenso wie die Größen $\varphi(h; q^t)$ der besonderen Reste φ , welche wir soeben betrachtet haben, in zwei Faktoren von wesentlich verschiedener Art zerfallen. Der eine dieser Faktoren, den man durch $\overset{0}{f}(h; q^t)$ bezeichnen kann, hängt, wenn $q=p$ ist, allein von den Zahlen p^{ω_k} , und wenn $q=2$ ist, allein von den Zahlen 2^{ω_k} und σ_k ab, während der zweite sich im Falle $q=p$ mittels einer, im Falle $q=2$ mittels einer oder zweier Einheiten C darstellen läßt. — Diese Einheiten C werden sich dann umgekehrt durch Größen $f(h; q^t)$ und $\overset{0}{f}(h; q^t)$ ausdrücken lassen. Da aber sowohl die Größen $f(h; q^t)$ als die Größen $\overset{0}{f}(h; q^t)$ für alle in bezug auf den Modul q^t kongruenten Formenklassen f gleiche Werte erhalten, so werden auch die Einheiten C für alle nach

dem Modul q^t kongruenten Formenklassen f die gleichen Werte annehmen, d. h. diese Einheiten werden Charaktere der Formen f vorstellen.

Die Charaktere C , zu welchen wir auf diese Weise gelangen, werden sich zunächst als Funktionen der Koeffizienten eines Hauptrestes $\varphi \pmod{q^t}$ darstellen. Nun haben wir aber in Kap. V gesehen, erstens, daß wir die Koeffizienten eines Hauptrestes φ durch die $n + 1$ Zahlen φ_h eines zum Hauptreste φ gehörigen Hauptrepräsentanten ausdrücken können, und zweitens, daß sich eine jede Grundform $\psi \pmod{q}$ einer Klasse f in einen Hauptrepräsentanten $\varphi \pmod{q^t}$ transformieren läßt, dessen Zahlen φ_h den Zahlen ψ_h gleich sind. Wir können demnach auch die Einheiten C durch die Zahlen φ_h eines Hauptrepräsentanten φ ausdrücken, und wir werden alsdann die Werte dieser Einheiten unmittelbar aus einer jeden beliebigen Grundform $\psi \pmod{q}$ herleiten können.

I. Es möge zunächst $q = p$ sein, und es sei f eine in bezug auf p primitive Form. Jeder Hauptrest φ der Klasse f für einen Modul p^t ($> p^{2n-1(p)}$) besitzt dann die Gestalt

$$\varphi \equiv \varphi_1 + \varphi_2 + \cdots + \varphi_\lambda \pmod{p^t},$$

$$\varphi_i \equiv p^{v_{\varphi_i-1}} \sum_{s=1}^{x_i} \alpha_s^{(i)} \xi_s^{(i)} \xi_s^{(i)} \pmod{p^t}.$$

Wir finden daher nach Satz J. in Kap. VII die Größen $\varphi(h; p^t) = f(h; p^t)$ durch die Gleichungen

$$(21) \quad \varphi(h; p^t) = \prod_{i=1}^{\lambda} \varphi_i(h; p^t).$$

Die $\lambda + 1$ Zahlen

$$M_0 = 0; \quad M_1 = v_{\varphi_1}, \quad M_2 = v_{\varphi_2}, \quad \dots, \quad M_{\lambda-1} = v_{\varphi_{\lambda-1}}; \quad M_\lambda = t$$

befriedigen die Ungleichungen

$$(22) \quad M_0 < M_1 < M_2 < \cdots < M_{\lambda-1} < M_\lambda.$$

Denn für ein $i < \lambda$ ist immer $M_i - M_{i-1} = v_{\varphi_i} - v_{\varphi_{i-1}} = \omega_{\varphi_i} \geq 1$, während wir für $i = \lambda$ gemäß unserer Annahme über t die Beziehung $M_\lambda - M_{\lambda-1} = t - v_{\varphi_{\lambda-1}} > 0$ bekommen.

Wir wollen die $\lambda + 1$ Zahlen

$$\bar{h}_0 = t - M_0 = t; \quad \bar{h}_1 = t - M_1; \quad \bar{h}_2 = t - M_2, \quad \dots, \quad \bar{h}_{\lambda-1} = t - M_{\lambda-1};$$

$$\bar{h}_\lambda = t - M_\lambda = 0$$

eingeführen. Aus den Ungleichungen (22) ergeben sich sofort die Ungleichungen

$$(23) \quad \bar{h}_0 > \bar{h}_1 > \bar{h}_2 > \cdots > \bar{h}_{\lambda-1} > \bar{h}_\lambda.$$

Wenn wir die Zahl h in der Form $h = p^{\bar{h}} \cdot h_0$ (h_0 relativ prim zu p) darstellen, so gelten für die Größen $\varphi_k(h; p^t)$ nach Kap. VIII, Absatz II, die Relationen

$$(24_0) \quad \text{für } \bar{h} \geq \bar{h}_{k-1}: \quad \varphi_k(h; p^t) = \overset{0}{\varphi}_k(h; p^t),$$

$$(24_1) \quad \text{für } \bar{h} < \bar{h}_{k-1}: \quad \varphi_k(h; p^t) = \overset{0}{\varphi}_k(h; p^t) \cdot \left(\frac{(\varphi_k)}{p}\right)^{\bar{h}_{k-1} - \bar{h}},$$

in welchen die Faktoren $\overset{0}{\varphi}_k(h; p^t)$ nicht-verschwindende und allein von den Zahlen $h, p^t; \omega(p)$ abhängige Zahlen bedeuten.

Aus der Gleichung (21) und den Formeln (24) erkennen wir, daß die Größen $\varphi(h; p^t)$ außer von den Zahlen $h, p^t; \omega(p)$ nur von den Einheiten $\left(\frac{(\varphi_k)}{p}\right)$ abhängen können. Wir wollen diese Einheiten umgekehrt durch die Größen $\varphi(h; p^t)$ und $\overset{0}{\varphi}_k(h; p^t)$ ausdrücken.

Wenn wir der Zahl h die Werte $1, 2, \dots, p^t$ erteilen, so durchläuft der Exponent $\bar{h}$ das Intervall $t, t-1, \dots, 0$. Beachten wir, daß $\bar{h}_0 = t$ und $\bar{h}_\lambda = 0$ ist, so leuchtet ein, daß dieses Intervall mit dem Intervalle

$$\bar{h}_0, \bar{h}_1, \bar{h}_2, \dots, \bar{h}_{\lambda-1}, \bar{h}_\lambda$$

übereinstimmt. Falls wir von der Größe $\bar{h} = t$ absehen, für welche offenbar $h = p^{\bar{h}} \cdot \bar{h}_0 \equiv 0 \pmod{p^t}$ und $\varphi(h; p^t) = p^{n^t}$ ist, wird daher eine jede der Zahlen $\bar{h} (= t-1, \dots, 0)$ einer und nur einer Ungleichung

$$(25) \quad \bar{h}_{i-1} > \bar{h} \geq \bar{h}_i$$

genügen. Wir wollen die Zahlen $\bar{h}$, welche der Ungleichung (25) genügen, durch $\bar{h}^{(i)}$ und die diesen Zahlen $\bar{h}^{(i)}$ zugehörigen Zahlen h durch $h^{(i)}$ bezeichnen; die Zahlen h , deren $\bar{h}$ gleich $\bar{h}_i$ ist, mögen h_i heißen.

Infolge (23) erfüllt jede Zahl $\bar{h}^{(i)}$ die Ungleichungen

$$(26_0) \quad \bar{h}^{(i)} \geq \bar{h}_i > \bar{h}_{i+1} > \dots > \bar{h}_{\lambda-1} > \bar{h}_\lambda,$$

$$(26_1) \quad \bar{h}^{(i)} < \bar{h}_{i-1} < \bar{h}_{i-2} < \dots < \bar{h}_1 < \bar{h}_0.$$

Folglich gehorchen alle Größen $\varphi_k(h^{(i)}; p^t)$, deren Index k einer der Zahlen $i+1, i+2, \dots, \lambda$ gleich ist, den Gesetzen (24₀) und alle Größen $\varphi_k(h^{(i)}; p^t)$, deren Index k einer der Zahlen $1, 2, \dots, i$ gleich ist, den Gesetzen (24₁). Für den Quotienten

$$\frac{\varphi_k(h^{(i)}; p^t)}{\varphi_k(h_{i-1}; p^t)} : \frac{\overset{0}{\varphi}_k(h^{(i)}; p^t)}{\overset{0}{\varphi}_k(h_{i-1}; p^t)} = \{h^{(i)}\}_k$$

gewinnen wir demnach die Formeln

$$\{h^{(i)}\}_k = 1 \quad \text{für } k = i+1, i+2, \dots, \lambda$$

und

$$\{h^{(i)}\}_k = \left(\frac{(\varphi_k)}{p}\right)^{\bar{h}_{i-1} - \bar{h}^{(i)}} \quad \text{für } k = 1, 2, \dots, i.$$

Setzen wir die Einheit

$$\left(\frac{\prod_{k=1}^i (\varphi_k)}{p} \right) = \Theta_i,$$

so ergibt sich für das Produkt $\prod_{k=1}^{\lambda} \varphi_k(h^{(i)}; p^t) = \varphi(h^{(i)}; p^t)$:

$$(27) \quad \varphi(h^{(i)}; p^t) = \Theta_i^{\bar{h}_{i-1} - \bar{h}^{(i)}} \cdot \varphi(h_{i-1}; p^t) \cdot \prod_{k=1}^{\lambda} \frac{\varphi_k(h^{(i)}; p^t)}{\varphi_k(h_{i-1}; p^t)}.$$

Nehmen wir jetzt an, es seien schon alle diejenigen Größen $\varphi(p^{\bar{h}} \cdot h_0; p^t)$ bekannt, deren $\bar{h}$ die Ungleichung $\bar{h} \geq \bar{h}_{i-1}$ befriedigt. Um alsdann zu einer Kenntnis aller Größen $\varphi(p^{\bar{h}} \cdot h_0; p^t)$ zu gelangen, deren $\bar{h} \geq \bar{h}_i$ ist, brauchen wir nur noch diejenigen Größen $\varphi(p^{\bar{h}} \cdot h_0; p^t)$ aufzusuchen, für welche $\bar{h}_{i-1} > \bar{h} \geq \bar{h}_i$ ist, das sind die Größen $\varphi(h^{(i)}; p^t)$. Aus der vorstehenden Gleichung (27) ist nun einerseits klar, daß wir zur Bestimmung dieser Größen höchstens der Einheiten $\Theta_i^{\bar{h}_{i-1} - \bar{h}^{(i)}}$ bedürfen, während andererseits diese Einheiten sich durch Größen $\varphi(h; p^t)$ und $\varphi_k(h; p^t)$ ausdrücken lassen. Unter den Einheiten $\Theta_i^{\bar{h}_{i-1} - \bar{h}^{(i)}}$ befindet sich, da die Differenz $\bar{h}_{i-1} - \bar{h}^{(i)}$ die Werte

$$1, 2, \dots, \bar{h}_{i-1} - \bar{h}_i \geq 1$$

durchläuft, die Einheit Θ_i selber. Wenden wir dieses Resultat der Reihe nach für die Werte $i = 1, 2, \dots, \lambda$ an, so erhalten wir den Satz:

Die Größen $\varphi(h; p^t)$ [$t > v_{n-1}(p)$] hängen außer von den Zahlen $p^{\omega(p)}$ von den $\lambda_p + 1$ Einheiten

$$\Theta_i = \left(\frac{\prod_{k=1}^i (\varphi_k)}{p} \right) \quad (i = 0, 1, 2, \dots, \lambda_p), \quad [\Theta_0 = 1]$$

ab, und umgekehrt können die $\lambda_p + 1$ Einheiten durch Größen $\varphi(h; p^t)$ und durch die Zahlen $p^{\omega(p)}$ ausgedrückt werden.

Sowohl die Einheiten Θ_i als die Einheiten $\left(\frac{(\varphi_i)}{p} \right) = \frac{\Theta_i}{\Theta_{i-1}}$ sind also Charaktere der Form f .

Wir drücken jetzt die Charaktere Θ_i durch die $n + 1$ Zahlen φ_h eines Hauptrepräsentanten aus, welcher den Rest φ liefert. Zunächst erkennen wir, daß die beiden Charaktere Θ_0 und Θ_λ nur von der Ordnung der Form f abhängen. Denn es ist

$$\Theta_0 = 1 = 1 \cdot \left(\frac{\varphi_0}{p} \right),$$

und für Θ_λ erhält man mit Hilfe der Kongruenz

$$\prod_{k=1}^{\lambda} (\varphi_k) \equiv (-1)^I \cdot \frac{d_{n-1}}{p^{\partial_{n-1}}} \pmod{p^{t-v_{n-1}}}, \quad [t > v_{n-1}]$$

die aus Kap. VI (S. 40) folgt, die Beziehung

$$\Theta_i = \left(\frac{(-1)^I \frac{d_{n-1}}{p^{\partial_{n-1}}}}{p} \right) = \left(\frac{\frac{d_{n-1}}{p^{\partial_{n-1}}}}{p} \right) \cdot \left(\frac{\varphi_n}{p} \right).$$

Die übrigen Charaktere Θ_i lassen sich mit Hilfe der Kongruenzen

$$\prod_{k=1}^i (\varphi_k) \equiv \varphi_{\mathfrak{g}_i} \frac{d_{\mathfrak{g}_i-1}}{p^{\partial_{\mathfrak{g}_i-1}}} \pmod{p^{\bar{h}_i-1}}$$

durch die Einheiten

$$C_i(p) = \left(\frac{\varphi_{\mathfrak{g}_i}}{p} \right)$$

ausdrücken. Diese Charaktere $C(p)$ sind mit den in Kap. VI gefundenen Charakteren $\left(\frac{\varphi_h}{p} \right)$ identisch.

Die Anzahl derjenigen Charaktere $C(p)$, welche nicht schon von vornherein durch die Ordnung der Form φ gegeben sind, beträgt $\lambda_p - 1$.

Nun ist für alle Primzahlen p , welche in dem Produkte $\prod_{k=1}^{n-1} o_k$ nicht aufgehen, $\lambda_p = 1$. Folglich werden nur diejenigen Primzahlen p , welche in diesem Produkt enthalten sind, besondere Charaktere $C(p)$ liefern. Auf diesem Umstande beruht es vornehmlich, daß die Anzahl aller unabhängigen Charaktere, welche einer gegebenen Formenklasse f angehören, einen endlichen Wert besitzt.

II. Es möge jetzt $q = 2$ sein, und es stelle f eine in bezug auf 2 primitive Form vor. Wir betrachten irgend einen Hauptrepräsentanten φ der Klasse f für einen Modul $2^t > \frac{4}{\sigma_{n-1}} \cdot 2^{v_{n-1}(2)}$. Sobald irgendeine Invariante σ_T der Form f den Wert 1 erhält, verschwinden von den Koeffizienten R_{ik} ($i < k$) des Restes φ alle diejenigen modulo 2^t , für welche $i \leq T$ und $k > T$ ist. Wenn $\sigma_T = 1$ ist, zerfällt demgemäß der Rest φ für den Modul 2^t in zwei Einzelreste Φ und Φ_0 , von denen der eine, Φ , die T ersten Reihen des Restes φ und der andere die $n - T$ letzten Reihen dieses Restes enthält. Wir überzeugen uns ferner, daß der Rest Φ zu 2 primitiv wird, während die höchste, in allen Koeffizienten des Restes Φ_0 aufgehende Potenz von 2 den Wert 2^{v_T} annimmt, sowie daß die Invarianten σ des Restes Φ gleich $\sigma_1, \sigma_2, \dots, \sigma_{T-1}$ und die Invarianten σ des Restes Φ_0 gleich $\sigma_{T+1}, \sigma_{T+2}, \dots, \sigma_{n-1}$ werden. Wiederholen wir diese Bemerkungen für mehrere Invarianten $\sigma_T = 1$, so gelangen wir ohne Schwierigkeit zu dem folgenden Satze:

K. Wenn wir von den $n - 1$ Invarianten σ der Form f beliebige $L - 1$ auswählen, welche den Wert 1 haben, etwa

$$(T_0 = 0) \quad \sigma_{T_1} = 1, \sigma_{T_2} = 1, \dots, \sigma_{T_{L-2}} = 1, \sigma_{T_{L-1}} = 1, \quad (T_L = n)$$

und die Zahlen $T_i - T_{i-1} = K_i$ setzen, so können wir aus den K_1 ersten Reihen von φ einen Rest Φ_1 , aus den K_2 folgenden Reihen einen Rest Φ_2 , usw., aus den K_L letzten Reihen von φ eine Form Φ_L bilden, und der Rest φ zerfällt für den Modul 2^t in L Einzelreste

$$\varphi \equiv \Phi_1 + \Phi_2 + \dots + \Phi_L \pmod{2^t}.$$

Die Reste Φ_i erscheinen als Produkte aus einem Faktor $2^{v_{T_i-1}}$ und einer zu 2 primitiven Form. Bezeichnen wir die $K_i - 1$ Invarianten σ dieser zu 2 primitiven Form durch $P_k^{(i)}$ und die $K_i - 1$ Invarianten $\omega(2)$ dieser Form durch $\Omega_k^{(i)}$, so bestehen die Relationen

$$P_k^{(i)} = \sigma_{T_{i-1}+k}, \quad \Omega_k^{(i)} = \omega_{T_{i-1}+k}. \quad (k = 1, 2, \dots, K_i - 1)$$

Einige besondere Fälle dieser Zerlegungen von φ in Formen Φ_i haben wir schon im Früheren untersucht.

So bekommen wir erstens die Zerlegung von φ in lauter Formen

$$F \equiv 2^M \alpha \xi^2 \pmod{2^t}$$

von einer Variablen oder Formen

$$F_0 \equiv 2^M (\alpha \xi^2 + \mathfrak{A} \xi \bar{\xi} + \bar{\alpha} \bar{\xi}^2) \pmod{2^t}$$

von zwei Variablen, indem wir für die Größen T_i ohne Ausnahme alle diejenigen Zahlen aus der Reihe 1, 2, ..., $n - 1$ wählen, denen eine Invariante $\sigma_T = 1$ entspricht.

Zweitens entsteht die in den Kapiteln III und IV betrachtete Zerlegung von φ in λ Formen φ_i , indem wir die Zahlen T_i gleich den $\lambda - 1$ Zahlen ϑ_i annehmen, welche zu den $\lambda - 1$ nicht-verschwindenden Zahlen ω_{ϑ_i} gehören.

Eine dritte Zerlegung des Restes φ , welche gleichfalls durch das angegebene Verfahren gefunden werden kann, wird uns jetzt zur Bestimmung der Größen $\varphi(h; 2^t)$ dienen.

Es mögen von den $n - 1$ Größen $\sigma_{i-1} 2^{w_i} \sigma_{i+1}$ im ganzen $\mu - 1$ ($1 \leq \mu \leq n$) durch 4 teilbar sein, nämlich die folgenden

$$(\eta_0 = 0) \quad \sigma_{\eta_1-1} 2^{w_{\eta_1}} \sigma_{\eta_1+1}, \dots, \sigma_{\eta_{\mu-1}-1} 2^{w_{\eta_{\mu-1}}} \sigma_{\eta_{\mu-1}+1}, \quad (\eta_{\mu} = n)$$

während alle übrigen $n - \mu$ dieser Zahlen entweder gleich 1 oder gleich 2 sein mögen. Nach Kap. IV, Absatz II bestehen alsdann die Relationen

$$\sigma_{\eta_1} = 1, \sigma_{\eta_2} = 1, \dots, \sigma_{\eta_{\mu-2}} = 1, \sigma_{\eta_{\mu-1}} = 1.$$

Wir können demnach die Zahlen T_i gleich den $\mu - 1$ Zahlen $\eta_1, \eta_2, \dots$,

$\eta_{\mu-2}, \eta_{\mu-1}$ annehmen, und wir gewinnen dann eine Zerlegung von φ in μ Formen

$$\varphi \equiv \Phi_1 + \Phi_2 + \dots + \Phi_\mu \equiv \sum_{i=1}^{\mu} \Phi_i \pmod{2^t}.$$

Die Größen $\varphi(h; 2^t) = f(h; 2^t)$ werden sich dann aus den Gleichungen

$$\varphi(h; 2^t) = \prod_{i=1}^{\mu} \Phi_i(h; 2^t)$$

bestimmen.

Nach dem Satze K. müssen die Zahlen

$$P_{k-1}^{(i)} 2^{\Omega_k^{(i)}} P_{k+1}^{(i)} \quad (k=1, 2, \dots, \eta_i - \eta_{i-1} - 1)$$

mit den Zahlen

$$\sigma_{\eta_{i-1}+k-1} 2^{\omega_{\eta_{i-1}+k}} \sigma_{\eta_{i-1}+k+1}$$

übereinstimmen. Da nun diese letzteren Zahlen unserer Annahme gemäß sämtlich kleiner als 4 sein sollen, so folgt, daß auch die Zahlen $P_{k-1}^{(i)} 2^{\Omega_k^{(i)}} P_{k+1}^{(i)}$ sämtlich gleich 1 oder gleich 2 sein werden. Wir überzeugen uns nun leicht mit Hilfe der in Kap. IV, Absatz II aufgestellten Sätze, daß dieses dann und nur dann eintritt, wenn die Reste Φ_i von der Gestalt

$$(28_1) \quad \Phi_i \equiv 2^{v_{\eta_{i-1}}} \sum_{s=1}^{m_i} \left(2^{s-1} \sum_{r=1}^{l_s} \alpha_r^{(s)} \xi_r^{(s)} \xi_r^{(s)} \right) \pmod{2^t}$$

$$[l_s \geq 1, m_i \geq 1]$$

oder von der Gestalt

$$(28_2) \quad \Phi_i \equiv 2^{v_{\eta_{i-1}+1}} (\alpha \xi^2 + \Re \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2) \pmod{2^t}$$

sind.

Wir wollen jeder Form Φ_i , je nachdem sie von der Gestalt (28₁) oder von der Gestalt (28₂) ist, eine Zahl τ_i gleich 1 oder gleich 2 zuteilen. Wir haben die Beziehungen

$$\tau_i = P_1^{(i)} = P_{\eta_i - \eta_{i-1} - 1}^{(i)}$$

oder

$$\tau_i = \sigma_{\eta_{i-1}+1} = \sigma_{\eta_i-1}.$$

Wir führen $\mu + 1$ Zahlen M_i mittels der Gleichungen

$$2^{M_0} = \tau_1 = \sigma_1, \quad 2^{M_1} = \tau_2 \cdot 2^{v_{\eta_1}}, \quad 2^{M_2} = \tau_3 \cdot 2^{v_{\eta_2}}, \quad \dots,$$

$$2^{M_{\mu-1}} = \tau_\mu \cdot 2^{v_{\eta_{\mu-1}}}, \quad 2^{M_\mu} = 2^t,$$

und μ Zahlen $\tilde{M}_i$ mittels der Gleichungen

$$2^{\tilde{M}_0} = \frac{4}{\tau_1} \cdot 2^{v_{\eta_1}-1}, \quad 2^{\tilde{M}_1} = \frac{4}{\tau_2} \cdot 2^{v_{\eta_2}-1}, \quad 2^{\tilde{M}_2} = \frac{4}{\tau_3} \cdot 2^{v_{\eta_3}-1}, \quad \dots,$$

$$2^{\tilde{M}_{\mu-1}} = \frac{4}{\tau_\mu} \cdot 2^{v_{\eta_\mu}-1}$$

ein. Die Potenzen 2^{M_i-1} werden, wenn $\tau_i = 1$ ist, gleich $2^{v_{\eta_i-1}}$, und wenn $\tau_i = 2$ ist, gleich $2^{v_{\eta_i-1}+1}$; die Potenzen $2^{\tilde{M}_i-1}$ sind, wenn $\tau_i = 1$ ist, gleich $2^{2+v_{\eta_i-1}} = 2^{1+m_i+M_i-1}$, dagegen wenn $\tau_i = 2$ ist, gleich

$$2^{1+v_{\eta_i-1}} = 2^{M_i-1}.$$

Folglich hat man stets $\tilde{M}_{i-1} \geq M_{i-1}$. Ferner erhält man wegen

$$\sigma_{\eta_i-1} 2^{\omega_{\eta_i}} \sigma_{\eta_i+1} \geq 4 \quad \text{und} \quad 2^t > \frac{4}{\sigma_{n-1}} \cdot 2^{v_{n-1}} \quad \text{für } i < \mu:$$

$$2^{M_i-\tilde{M}_{i-1}} = \frac{\tau_i \tau_{i+1} \cdot 2^{v_{\eta_i}-v_{\eta_i-1}}}{4} = \frac{\sigma_{\eta_i-1} \cdot 2^{\omega_{\eta_i}} \cdot \sigma_{\eta_i+1}}{4} \geq 1$$

und

$$2^{M_\mu-\tilde{M}_{\mu-1}} = \frac{2^t}{\frac{4}{\sigma_{n-1}} \cdot 2^{v_{n-1}}} > 1,$$

so daß man die Ungleichungen findet

$$(29) \quad M_0 \leq \tilde{M}_0 \leq M_1 \leq \tilde{M}_1 \leq \dots \leq M_{\mu-1} \leq \tilde{M}_{\mu-1} < M_\mu = t.$$

Wir bezeichnen die Zahlen $t - M_i$ durch $\bar{h}_i$ und die Zahlen $t - \tilde{M}_i$ durch $\tilde{h}_i$. Aus (29) folgen sofort die Ungleichungen

$$(30) \quad \bar{h}_0 \geq \tilde{h}_0 \geq \bar{h}_1 \geq \tilde{h}_1 \geq \dots \geq \bar{h}_{\mu-1} \geq \tilde{h}_{\mu-1} > \bar{h}_\mu = 0,$$

und es gelten die Relationen

$$\text{für } \tau_i = 1: \quad \bar{h}_{i-1} = \tilde{h}_{i-1} + 1 + m_i,$$

$$\text{für } \tau_i = 2: \quad \bar{h}_{i-1} = \tilde{h}_{i-1},$$

und

$$(31) \quad 2^{\tilde{h}_{i-1}-\bar{h}_i} = \frac{\sigma_{\eta_i-1} \cdot 2^{\omega_{\eta_i}} \cdot \sigma_{\eta_i+1}}{4} \geq 1 \quad (i < \mu);$$

$$\tilde{h}_{\mu-1} - \bar{h}_\mu > 0.$$

Wir wollen für einen jeden Rest Φ_i von der Gestalt (28₁)

$$\Phi_i \equiv 2^{M_i-1} \sum_{s=1}^{m_i} \left(2^{s-1} \sum_{r=1}^{l_s} \alpha_r^{(s)} \xi_r^{(s)} \xi_r^{(s)} \right) \equiv \sum_{s=1}^{m_i} \varphi_s^{(i)} \pmod{2^t}$$

eine Einheit

$$1_i = \left(\frac{2}{\left(\varphi_{m_i-1}^{(i)} \right)} \right) \left(\frac{2}{\left(\varphi_{m_i-3}^{(i)} \right)} \right) \cdots \left(\frac{2}{\left(\varphi_{m_i+1-2\left[\frac{m_i}{2}\right]}^{(i)} \right)} \right) \cdot \left(\frac{2}{\prod_{k=1}^{i-1} (\Phi_k)} \right)^{m_i+1-2\left[\frac{m_i}{2}\right]} \cdot (-1)^{\frac{1}{2} \sum_{(r',s') \neq (r'',s'')} \frac{\alpha_{r'}^{s'}-1}{2} \cdot \frac{\alpha_{r''}^{s''}-1}{2}}$$

und für einen jeden Rest Φ_i von der Gestalt (28₂)

$$\Phi_i \equiv 2^{M_i-1}(\alpha \xi^2 + \mathfrak{A} \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2) \pmod{2^f}$$

eine Einheit

$$l_i = 1$$

bilden.

Für die Größen $\Phi_k(h; 2^f)$ [$h = 2^{\bar{h}} h_0$, $h_0 \equiv 1 \pmod{2}$] ergeben sich jetzt die Formeln

$$(32_0) \quad \text{für } \bar{h} \geq \bar{h}_{k-1}: \quad \Phi_k(h; 2^f) = \overset{0}{\Phi}_k(h; 2^f),$$

$$(32_1) \quad \text{für } \bar{h} \leq \tilde{h}_{k-1}: \quad \Phi_k(h; 2^f) = \overset{0}{\Phi}_k(h; 2^f) \cdot l_k \cdot \left(\frac{2}{k-1} \right)^{\bar{h}_{k-1} - \tilde{h}_{k-1}} \cdot \prod_{i=1}^{(\Phi_k)} \cdot (-i)^{\tau_k^2 \left(\frac{(\Phi_k)-1}{2} \right)^2} \cdot (-1)^{\tau_k \cdot \frac{(\Phi_k)-1}{2} \cdot \frac{h_0-1}{2} \left(\frac{2}{(\Phi_k)} \right)^{\tilde{h}_{k-1} - \bar{h}}},$$

$$(32_2) \quad \text{und für } \tau_k = 1, \bar{h}_{k-1} > \bar{h} > \tilde{h}_{k-1}: \quad \Phi_k(h; 2^f) = 0,$$

in welchen die Faktoren $\overset{0}{\Phi}_k(h; 2^f)$ nicht-verschwindende und allein von den Zahlen $h, 2^f$; $\sigma, \omega(2)$ abhängige Größen bedeuten.

Wir wollen die Zahlen $\bar{h}$, welche der Ungleichung

$$\tilde{h}_{k-1} \geq \bar{h} \geq \bar{h}_k$$

genügen, durch $\bar{h}^{(k)}$ bezeichnen. Eine jede Zahl $\bar{h}^{(k)}$ befriedigt infolge (30) außerdem die Ungleichungen

$$(33_1) \quad \bar{h}^{(k)} \geq \bar{h}_k \geq \bar{h}_{k+1} \geq \dots \geq \bar{h}_{\mu-1} \geq \bar{h}_\mu,$$

$$(33_2) \quad \bar{h}^{(k)} \leq \tilde{h}_{k-1} \leq \tilde{h}_{k-2} \leq \dots \leq \tilde{h}_1 \leq \tilde{h}_0.$$

Demzufolge gehorchen alle Größen $\Phi_i(2^{\bar{h}^{(k)}} h_0; 2^f)$, deren Index i einer der Zahlen $k+1, k+2, \dots, \mu$ gleich ist, den Gesetzen (32₀) und alle Größen $\Phi_i(2^{\bar{h}^{(k)}} h_0; 2^f)$, deren Index i einer der Zahlen $1, 2, \dots, k$ gleich ist, den Gesetzen (32₁). Wir bekommen demnach für den Quotienten

$$\frac{\Phi_i(2^{\bar{h}^{(k)}} h_0; 2^f)}{\Phi_i(2^{\tilde{h}_{k-1}} h_0; 2^f)} : \frac{\overset{0}{\Phi}_i(2^{\bar{h}^{(k)}} h_0; 2^f)}{\overset{0}{\Phi}_i(2^{\tilde{h}_{k-1}} h_0; 2^f)} = \{h^{(k)}\}_i$$

die Gleichungen

$$\{h^{(k)}\}_i = 1 \quad (\text{für } i = k+1, k+2, \dots, \mu),$$

$$\{h^{(k)}\}_i = \left(\frac{2}{(\Phi_i)} \right)^{\bar{h}_{k-1} - \bar{h}^{(k)}} \quad (\text{für } i = 1, 2, \dots, k-1)$$

und für $i = k$:

$$\{h^{(k)}\}_k = l_k \cdot \left(\frac{2}{k-1} \right)^{\bar{h}_{k-1} - \tilde{h}_{k-1}} \cdot (-i)^{\tau_k^2 \left(\frac{(\Phi_k)-1}{2} \right)^2} \cdot (-1)^{\tau_k \frac{(\Phi_k)-1}{2} \cdot \frac{h_0-1}{2} \left(\frac{2}{(\Phi_k)} \right)^{\tilde{h}_{k-1} - \bar{h}^{(k)}}}.$$

Für das Produkt $\prod_{i=1}^{\mu} \Phi_i(2^{\bar{h}^{(k)}} h_0; 2^f) = \varphi(2^{\bar{h}^{(k)}} h_0; 2^f)$ ergibt sich jetzt

$$\varphi(2^{\bar{h}(k)} h_0; 2^t) = l_k \cdot (-i)^{\tau_k \frac{(\Phi_k)-1}{2}} (-1)^{\tau_k \frac{(\Phi_k)-1}{2} \cdot \frac{h_0-1}{2}} \cdot \left(\frac{2}{\prod_{i=1}^{\mu} (\Phi_i)} \right)^{\tilde{h}_{k-1} - \bar{h}(k)} \cdot \varphi(2^{\tilde{h}_{k-1}} h_0; 2^t) \prod_{i=1}^{\mu} \frac{\Phi_i(2^{\bar{h}(k)} h_0; 2^t)}{\Phi_i(2^{\tilde{h}_{k-1}} h_0; 2^t)},$$

woraus wir insbesondere die Beziehungen

$$(34_1) \quad \varphi(2^{\tilde{h}_{k-1}} h_0; 2^t) = l_k \cdot (-i)^{\tau_k \frac{(\Phi_k)-1}{2}} (-1)^{\tau_k \frac{(\Phi_k)-1}{2} \cdot \frac{h_0-1}{2}} \cdot \varphi(2^{\bar{h}_{k-1}} h_0; 2^t) \prod_{i=1}^{\mu} \frac{\Phi_i(2^{\tilde{h}_{k-1}} h_0; 2^t)}{\Phi_i(2^{\bar{h}_{k-1}} h_0; 2^t)}$$

und

$$(34_2) \quad \varphi(2^{\bar{h}(k)} h_0; 2^t) = \left(\frac{2}{\prod_{i=1}^{\mu} (\Phi_i)} \right)^{\tilde{h}_{k-1} - \bar{h}(k)} \varphi(2^{\tilde{h}_{k-1}} h_0; 2^t) \prod_{i=1}^{\mu} \frac{\Phi_i(2^{\bar{h}(k)} h_0; 2^t)}{\Phi_i(2^{\tilde{h}_{k-1}} h_0; 2^t)}$$

gewinnen.

Wir nehmen jetzt an, es seien schon alle diejenigen Größen $\varphi(2^{\bar{h}} h_0; 2^t)$ bekannt, deren $\bar{h} \geq \bar{h}_{k-1}$ ist. Um alsdann zu einer Kenntnis aller Größen $\varphi(2^{\bar{h}} h_0; 2^t)$ zu gelangen, deren $\bar{h}$ die Ungleichung $\bar{h} \geq \bar{h}_k$ befriedigt, brauchen wir nur die sämtlichen Größen $\varphi(2^{\bar{h}} h_0; 2^t)$ zu untersuchen, für welche $\tilde{h}_{k-1} \geq \bar{h} \geq \bar{h}_k$ gilt, das sind die Größen $\varphi(2^{\bar{h}(k)} h_0; 2^t)$. Denn ist $\tau_k = 1$, so wird $\bar{h}_{k-1} = \tilde{h}_{k-1} + 1 + m_k$, und die sämtlichen Zahlen $\bar{h}$, welche $\geq \bar{h}_k$ sind, genügen entweder der Ungleichung $\bar{h} \geq \bar{h}_{k-1}$ oder der Ungleichung $\bar{h}_{k-1} > \bar{h} > \tilde{h}_{k-1}$ oder der Ungleichung $\tilde{h}_{k-1} \geq \bar{h} \geq \bar{h}_k$. Ist aber $\bar{h} \geq \bar{h}_{k-1}$, so ist die Größe $\varphi(2^{\bar{h}} h_0; 2^t)$ schon bekannt, und erfüllt $\bar{h}$ die Ungleichung $\bar{h}_{k-1} > \bar{h} > \tilde{h}_{k-1}$, so wird $\Phi_k(2^{\bar{h}} h_0; 2^t) = 0$, also auch $\varphi(2^{\bar{h}} h_0; 2^t) = 0$. Es sind mithin in der Tat nur diejenigen Größen $\varphi(2^{\bar{h}} h_0; 2^t)$ zu betrachten, für welche $\tilde{h}_{k-1} \geq \bar{h} \geq \bar{h}_k$ ist. Wenn hingegen $\tau_k = 2$ ist, so wird $\bar{h}_{k-1} = \tilde{h}_{k-1}$, und die Zahlen $\bar{h}$, welche $\geq \bar{h}_k$ sind, genügen entweder der Ungleichung $\bar{h} \geq \bar{h}_{k-1}$, in welchem Falle die Größen $\varphi(2^{\bar{h}} h_0; 2^t)$ als bekannt anzusehen sind, oder sie genügen der Ungleichung $\tilde{h}_{k-1} \geq \bar{h} \geq \bar{h}_k$.

Mit Hilfe der Gleichungen (34₁) und (34₂) ziehen wir jetzt den folgenden Schluß:

L. Wenn alle Größen $\varphi(2^{\bar{h}} h_0; 2^t)$, deren $\bar{h} \geq \bar{h}_{k-1}$ ist, gegeben sind, so bedürfen wir zur Bestimmung der sämtlichen Größen $\varphi(2^{\bar{h}(k)} h_0; 2^t)$ höchstens der Einheiten $l_k, (-1)^{\frac{(\Phi_k)-1}{2}}, \left(\frac{2}{\prod_{i=1}^{\mu} (\Phi_i)} \right)^{\tilde{h}_{k-1} - \bar{h}(k)}$, und umgekehrt

können diese Einheiten stets durch Größen $\varphi(2^{\bar{h}} h_0; 2^t)$ und $\Phi_i(2^{\bar{h}} h_0; 2^t)$ ($\bar{h} \geq \bar{h}_k$) ausgedrückt werden.

Die GröÙe $\tilde{h}_{k-1} - \bar{h}^{(k)}$ erhält, da die Zahlen $\bar{h}^{(k)}$ der Ungleichung $\tilde{h}_{k-1} \geq \bar{h}^{(k)} \geq \bar{h}_k$ genügen sollen, die Werte

$$0, 1, 2, \dots, \tilde{h}_{k-1} - \bar{h}_k.$$

Unter den Einheiten $\left(\frac{2}{\prod_{i=1}^k (\Phi_i)} \right)^{\tilde{h}_{k-1} - \bar{h}^{(k)}}$ wird daher die Einheit $\left(\frac{2}{\prod_{i=1}^k (\Phi_i)} \right)$

selber auftreten, sobald $\tilde{h}_{k-1} - \bar{h}_k \geq 1$ ist.

Wir wollen annehmen, von den Zahlen $\sigma_{i-1} 2^{\omega_i} \sigma_{i+1}$ ($i = 1, 2, \dots, n-2, n-1$) seien im ganzen $\nu-1$ ($1 \leq \nu \leq n$) größer oder gleich 8, nämlich die folgenden:

$$(\xi_0 = 0) \quad \sigma_{\xi_1-1} 2^{\omega_{\xi_1}} \sigma_{\xi_1+1}, \dots, \sigma_{\xi_{\nu-1}-1} 2^{\omega_{\xi_{\nu-1}}} \sigma_{\xi_{\nu-1}+1}, \quad (\xi_{\nu} = n),$$

wo

$$(K_0 = 0) \quad \xi_1 = \eta_{K_1}, \dots, \xi_{\nu-1} = \eta_{K_{\nu-1}} \quad (K_{\nu} = \mu)$$

ist, während alle übrigen $n-\nu$ Zahlen $\sigma_{i-1} 2^{\omega_i} \sigma_{i+1}$ höchstens gleich 4 sein mögen. Es sind dann infolge der Gleichungen (31) von den μ Zahlen $\tilde{h}_{k-1} - \bar{h}_k$ ($k = 1, 2, \dots, \mu$) die ν folgenden

$\tilde{h}_{K_1-1} - \bar{h}_{K_1}, \tilde{h}_{K_2-1} - \bar{h}_{K_2}, \dots, \tilde{h}_{K_{\nu-1}-1} - \bar{h}_{K_{\nu-1}}$ und $\tilde{h}_{K_{\nu}-1} - \bar{h}_{K_{\nu}}$ größer oder gleich 1, während die übrigen $\mu-\nu$ dieser Zahlen verschwinden.

Wenden wir das Resultat L. der Reihe nach für die Werte $k = 1, 2, \dots, \mu$ an und beachten wir einerseits, daß alle GröÙen $\varphi(2^{\bar{h}_0}; 2^t)$, deren $\bar{h} \geq \bar{h}_0$ ist, gleich 2^{nt} werden, und andererseits, daß ein jedes $\bar{h} \geq \bar{h}_{\mu}$, d. i. ≥ 0 ist, so gewinnen wir den Satz:

Die GröÙen $\varphi(h; 2^t)$ ($2^t > \frac{4}{\sigma_{n-1}} 2^{v_{n-1}(2)}$) hängen von den Zahlen σ , $2^{\omega(2)}$ und den $2(\mu+1) + (\nu+1)$ Einheiten

$$H_k = (-1)^{\frac{-1 + \prod_{i=1}^k (\Phi_i)}{2}} \quad (k = 0, 1, 2, \dots, \mu) \quad [H_0 = 1],$$

$$Z_k = \left(\frac{2}{\prod_{i=1}^k (\Phi_i)} \right) \quad (k = 0, 1, 2, \dots, \nu) \quad [Z_0 = 1],$$

und

$$E_k = \prod_{i=1}^k \left\{ \left(\frac{2}{\prod_{j=1}^{i-1} (\Phi_j)} \right)^{\tilde{h}_i - 2 - \bar{h}_{i-1}} \cdot 1_i \cdot (-1)^{\frac{-1 + \prod_{j=1}^{i-1} (\Phi_j)}{2} \cdot \frac{1 + \prod_{j=1}^i (\Phi_j)}{2}} \right\} \\ (k = 0, 1, 2, \dots, \mu) \quad [E_0 = 1]$$

ab, und umgekehrt können diese sämtlichen Einheiten durch die Zahlen σ , $2^{\omega(2)}$ und durch Größen $\varphi(h; 2^t)$ ausgedrückt werden.

Die Einheiten E_k, H_k, Z_k sind also ebenso wie die Einheiten l_k , $(-1)^{\frac{(\Phi_k)-1}{2}}$ Charaktere der Form f .

Die Anzahl der Reste Φ_k , denen ein $\tau_k = 2$ entspricht, ist gleich der Anzahl aller derjenigen Invarianten σ der Form f , welche gleich 2 werden.

Die Charaktere $(-1)^{\frac{(\Phi_k)-1}{2}}$, welche einer Form Φ_k von der Gestalt (28_2) angehören, sind gleich -1 , und die Charaktere l_k , welche einer Form Φ_k von der Gestalt (28_2) entsprechen, gleich $+1$. Ferner erkennen wir aus der am Schlusse von Kap. VIII gemachten Bemerkung, daß die Einheit l_k den Wert $+1$ erhalten muß, so oft der Rest Φ_k die Gestalt

$$\Phi_k \equiv 2^M \alpha \xi^2 \pmod{2^t}$$

oder die Gestalt

$\Phi_k \equiv 2^M (\alpha \xi^2 + \alpha_0 \xi_0^2) \pmod{2^t}$; $(-1)^{\frac{\alpha \alpha_0 - 1}{2}} = (-1)^{\frac{(\Phi_k) - 1}{2}} = -1$ besitzt.

Wir drücken jetzt die Charaktere E_k, H_k, Z_k durch die $n+1$ Zahlen φ_h eines Hauptrepräsentanten $\varphi \pmod{2^t}$ aus.

Die in Kap. VI (S. 40) aufgestellte Kongruenz ergibt unmittelbar

$$\prod_{i=1}^{\mu} (\Phi_i) \equiv (-1)^r \cdot \frac{d_{n-1}}{2^{\partial_{n-1}}} \left(\pmod{4 \cdot \frac{2^t}{\frac{4}{\sigma_{n-1}} 2^{v_{n-1}}}} \right);$$

auf einem ähnlichen Wege gelangen wir zu den Beziehungen

$$\prod_{i=1}^k (\Phi_i) \equiv \frac{d_{\eta_k-1}}{2^{\partial_{\eta_k-1}}} \cdot \varphi_{\eta_k} \pmod{2^{2+\tilde{h}_k-1}}.$$

Also wird

$$H_0 = (-1)^{\frac{\varphi_0-1}{2}}, \quad H_\mu = \left(\frac{-1}{\frac{d_{n-1}}{2^{\partial_{n-1}}}} \right) (-1)^{\frac{\varphi_n-1}{2}},$$

$$Z_0 = \left(\frac{2}{\varphi_0} \right), \quad Z_r = \left(\frac{2}{\frac{d_{n-1}}{2^{\partial_{n-1}}}} \right) \left(\frac{2}{\varphi_n} \right),$$

und ebenso kann man die übrigen Charaktere H_k, Z_k durch die Einheiten

$$C_k(4) = (-1)^{\frac{\varphi_{\eta_k}-1}{2}}$$

und

$$C_k(8) = \left(\frac{2}{\varphi_{\eta_k}} \right)$$

ausdrücken. Diese Charaktere $C(4)$ und $C(8)$ sind mit den in Kap. VI

betrachteten Charakteren $(-1)^{\frac{\varphi_h - 1}{2}}$ und $\left(\frac{2}{\varphi_h}\right)$ identisch. Die Charaktere E_k lassen sich als Produkte von zwei Faktoren darstellen, von denen der eine [falls $o_h = 2^{\omega_h} \cdot e_h$, $e_h \equiv 1 \pmod{2}$] gesetzt wird] gleich

$$\mathfrak{G}_k(4) = \prod_{i=1}^{\eta_k-1} \left(\frac{\sigma_{i-1} 2^{\omega_i} \sigma_{i+1}}{\varphi_i} \right) (-1)^{\sum_{i=1}^{\eta_k-1} \frac{\varphi_i-1}{2} \cdot \frac{e_i+1}{2}} \\ \cdot (-1)^{\frac{\varphi_0-1}{2} \cdot \frac{\varphi_1-1}{2} + \frac{\varphi_1-1}{2} \cdot \frac{\varphi_2-1}{2} + \dots + \frac{\varphi_{\eta_k-2}-1}{2} \cdot \frac{\varphi_{\eta_k-1}-1}{2} + \frac{\varphi_{\eta_k-1}-1}{2} \cdot \frac{\varphi_{\eta_k}-1}{2}}$$

ist, während der andere allein von der Ordnung der Form f und dem Charakter

$$C_k(4) = (-1)^{\frac{\varphi_{\eta_k}-1}{2}}$$

abhängt. Bei einer Aufzählung sämtlicher Charaktere für den Modul 2^t können wir daher die Charaktere E_k durch die Charaktere $\mathfrak{G}_k(4)$ ersetzen.

Kap. X. [[Bedingungen für die Gültigkeit der Kongruenz $f \equiv g \pmod{q^t}$.]]

Wir haben gesehen, daß die Kongruenz zweier Klassen f und g nach einem Modul q^t , der die höchsten in den Größen $\frac{4}{\sigma_{n-1}} o_1 o_2 \dots o_{n-1}$ der beiden Formen f und g aufgehenden Potenzen von q übersteigt, von den folgenden Bedingungen abhängt:

I. Die Klassen f und g besitzen, wenn $q = p$, dieselben Größen p^{ω_h} und, wenn $q = 2$, dieselben Größen 2^{ω_h} , σ_h .

II. Die Determinanten der Klassen f und g sind, wenn $q = p$ ist, nach dem Modul $p^{t+\partial_{n-2}(p)}$, und wenn $q = 2$ ist, nach dem Modul $\sigma_{n-1} \cdot 2^{t+\partial_{n-2}(2)}$ einander kongruent.

III. Die Hauptreste der Klassen f und g besitzen, wenn $q = p$, dieselben Charaktere $\Theta(p)$ und, wenn $q = 2$, dieselben Charaktere $E(4)$, $H(4)$, $Z(8)$; oder (was auf dasselbe hinauskommt):

Die Größen $f(h; q^t)$ und $g(h; q^t)$ sind identisch.

Diese Bedingungen sind aber auch hinreichend dafür, daß die Klassen f und g nach dem Modul q^t kongruent sind; denn es gilt der Satz:

M. Genügen zwei Klassen f und g den Bedingungen I, II, III, so kann man jede Form der einen in äquivalente Formen transformieren, welche nach dem Modul q^t einem beliebigen Repräsentanten g der andern kongruent sind.

Man kann sich zweier verschiedener Methoden bedienen, um diesen Satz zu beweisen, indem man ihn entweder durch einen Schluß von

$n_0 (< n)$ auf n bestätigt oder indem man Substitutionen bestimmt, welche einen Hauptrest φ der Klasse f in einen Hauptrest ψ der Klasse g überführen. — Da die Zeit drängt, muß ich darauf verzichten, diese Methoden hier zu entwickeln.*)

Kap. XI. Genera von Formen. — Bedingungen für die Existenz eines Genus.

Die Zahl der Charaktere, welche zu einer gegebenen Form f gehören und nicht von vornherein durch die Ordnung

$$O: \begin{pmatrix} \sigma_h \\ o_h \end{pmatrix}, \quad I$$

dieser Form bestimmt sind, ist stets endlich. Denn wir sahen, daß allein die in dem Produkt $2 \cdot \frac{2}{\sigma_{n-1}} o_1 o_2 \dots o_{n-1} = \Pi(f)$ enthaltenen Primzahlen derartige Charaktere liefern können.

Wir fassen alle diejenigen Formenklassen, welche dieselbe Ordnung und dieselben Charaktere wie eine gegebene Form besitzen, in ein *Genus* von Formen zusammen.

Der Satz M. in Kap. X zeigt uns, daß die Klassen eines und desselben Genus in bezug auf jeden Modul $p^t > p^{e_{n-1}(p)}$ und in bezug auf jeden Modul $2^t > \frac{4}{\sigma_{n-1}} 2^{e_{n-1}(2)}$ kongruent sind, woraus wir mit Hilfe des Satzes II in Kap. VI schließen, daß *zwei Klassen desselben Genus in bezug auf jeden beliebigen Modul kongruent sind.*

Umgekehrt ist klar, daß zwei Klassen f und g von gleichem Index I , welche verschiedenen Genera angehören, nicht in bezug auf jeden Modul N kongruent sein können. Denn wenn die Invarianten oder die Charaktere der Klassen f und g nicht die gleichen sind, so kann, wie man leicht erkennt, nicht gleichzeitig $f \equiv g \pmod{\Pi(f)}$ und $f \equiv g \pmod{\Pi(g)}$ sein.

Man kann die sämtlichen Charaktere einer Klasse f mit Hilfe eines Hauptrestes dieser Klasse für den Modul Π herleiten. Folglich werden zwei Klassen derselben Ordnung O die gleichen Charaktere besitzen, sobald sie in bezug auf den Modul Π einander kongruent sind. Also kann man ein Genus von Formen auch in der folgenden Weise definieren:

Ein Genus f besteht aus allen denjenigen Klassen, welche der Ordnung f angehören und in bezug auf den Modul $\Pi(f)$ der Klasse f kongruent sind.

I. Um zu entscheiden, ob zwei Klassen derselben Ordnung O in demselben Genus enthalten sind, kann man sich der Grundformen dieser Klassen für den Modul Π bedienen. — Wir behaupten:

*) Siehe die Note [[am Schlusse dieser Abhandlung, S. 136--143]].

Jede primitive Formenklasse f der Ordnung O besitzt für den Modul Π Grundformen φ , in welchen die Zahlen φ_h zu den Zahlen $\varphi_{h-1} \cdot \varphi_{h+1}$ relativ prim sind.

In der Tat: wir bemerken zunächst, daß eine Form $f_{h+1} = \{r_{ik}^{(h+1)}\}$ ($i, k = 1, 2, \dots, h+1$) von $h+1$ Variablen und von einer Ordnung

$$\begin{pmatrix} \sigma_1, \sigma_2, \dots, \sigma_{h-1}, \sigma_h \\ o_1, o_2, \dots, o_{h-1}, \sigma_{h+1} \varphi_{h+1} \cdot o_h \end{pmatrix}$$

stets in eine äquivalente Form $\{r_{ik}^{(h)}\}$ ($i, k = 1, 2, \dots, h+1$) transformiert werden kann, in welcher die Teilform $f_h = \{r_{ik}^{(h)}\}$ ($i, k = 1, 2, \dots, h$) von h Variablen eine Ordnung

$$\begin{pmatrix} \sigma_1, \sigma_2, \dots, \sigma_{h-2}, \sigma_{h-1} \\ o_1, o_2, \dots, o_{h-2}, \sigma_h \varphi_h \cdot o_{h-1} \end{pmatrix}$$

mit einer zu der Zahl $\Pi \cdot \varphi_{h+1}$ relativ primen Zahl φ_h besitzt. Um eine Form von der gewünschten Beschaffenheit zu erhalten, brauchen wir nämlich nur eine Grundform der Klasse f_{h+1} in bezug auf den Modul $\Pi \cdot \varphi_{h+1}$ zu bestimmen.

Wenden wir diese Bemerkung für $h = n-1, n-2, \dots, 1$ an, indem wir von der Form $f = f_n = \{r_{ik}^{(n)}\}$ ($i, k = 1, 2, \dots, n$) ausgehen, so gewinnen wir einen Repräsentanten $\varphi = \{r_{ik}^{(1)}\}$ ($i, k = 1, 2, \dots, n$), in welchem die Teilformen $\{r_{ik}^{(1)}\}$ ($i, k = 1, 2, \dots, h$) eine Ordnung

$$\begin{pmatrix} \sigma_1, \sigma_2, \dots, \sigma_{h-2}, \sigma_{h-1} \\ o_1, o_2, \dots, o_{h-2}, \sigma_h \varphi_h \cdot o_{h-1} \end{pmatrix}$$

besitzen, während die Zahlen φ_h zu den Zahlen $\Pi \cdot \varphi_{h+1}$ relativ prim sind. Dieser Repräsentant φ gibt uns also eine Grundform für den Modul Π mit Zahlen φ_h , welche zu den Zahlen $\varphi_{h-1} \cdot \varphi_{h+1}$ relativ prim sind. Eine solche Form φ nennen wir eine *charakteristische Form* der Klasse f .

Es möge irgendeine charakteristische Form φ der Klasse f vorliegen. Wir wollen die Formen von h Variablen, welche aus den h ersten Reihen des Systems von φ gebildet sind, durch $\{\varphi_h\}$ bezeichnen. Es mögen die Indizes dieser Formen $\{\varphi_h\}$ ($h = 1, 2, \dots, n$) gleich I_h sein. Dann ist $I_n = I$, und wir setzen noch $I_0 = 0$. Wir schreiben $(-1)^{I_h} = \varepsilon_h$. Die $n+1$ Einheiten ε_h stellen die Vorzeichen der Größen φ_h vor, und es werden daher die Größen $\varepsilon_h \varphi_h$ sämtlich positiv.

Für die Formen $\{\varphi_h\}$ erhalten wir nach einem bekannten Satz die Zerlegungen

$$\{\varphi_h\} = \frac{Z_1^2}{\varphi_0 \varphi_1} + \frac{Z_2^2}{\varphi_1 \varphi_2} + \dots + \frac{Z_h^2}{\varphi_{h-1} \varphi_h}.$$

Führen wir jetzt n Einheiten $\delta_1, \delta_2, \dots, \delta_n$ mittels der Beziehungen $\delta_h = \varepsilon_h \varepsilon_{h-1}$, $\varepsilon_h = \delta_1 \delta_2 \dots \delta_h$ ein, so müssen von den h ersten Einheiten

$\delta_1, \delta_2, \dots, \delta_h$ im ganzen I_h gleich -1 und $h - I_h$ gleich $+1$ sein, und wir bekommen insbesondere

$$(-1)^{\sum_{i < k} \frac{\delta_i - 1}{2} \cdot \frac{\delta_k - 1}{2}} = (-1)^{\frac{I_h(I_h - 1)}{2}} = (-1)^{\left[\frac{I_h}{2}\right]}.$$

Wenn wir diese Einheit durch die Größen ε_i ausdrücken, so ergibt sich

$$(-1)^{\frac{I_h(I_h - 1)}{2}} = (-1)^{\sum_{i=1}^{h-1} \frac{\varepsilon_i - 1}{2}} \cdot (-1)^{\frac{\varepsilon_0 - 1}{2} \cdot \frac{\varepsilon_1 - 1}{2} + \frac{\varepsilon_1 - 1}{2} \cdot \frac{\varepsilon_2 - 1}{2} + \dots + \frac{\varepsilon_{h-2} - 1}{2} \cdot \frac{\varepsilon_{h-1} - 1}{2} + \frac{\varepsilon_{h-1} - 1}{2} \cdot \frac{\varepsilon_h - 1}{2}}.$$

Aus der in Kap. VI (S. 43) aufgestellten Gleichung (8) gewinnen wir eine Kongruenz

$$(35) \quad -\sigma_{h-1} 2^{\omega_h} \sigma_{h+1} e_h \varphi_{h-1} \varphi_{h+1} \equiv X_h^2 \pmod{\sigma_h^2 \varphi_h},$$

in welcher die Zahl X_h zu $\sigma_h^2 \varphi_h$ relativ prim ausfällt, da nach unserer Voraussetzung über die Form φ die Zahl $-\sigma_{h-1} \sigma_h \sigma_{h+1} \varphi_{h-1} \varphi_{h+1}$ zu $\sigma_h^2 \varphi_h$ relativ prim ist. Aus dieser Kongruenz folgt auf der Stelle die Gleichung

$$\left(\frac{\varphi_{h-1} \varphi_{h+1}}{\varepsilon_h \varphi_h} \right) = \left(\frac{-\sigma_{h-1} 2^{\omega_h} \sigma_{h+1}}{\varepsilon_h \varphi_h} \right) \cdot \left(\frac{e_h}{\varepsilon_h \varphi_h} \right),$$

welche wegen

$$\left(\frac{-1}{\varepsilon_h \varphi_h} \right) = (-1)^{\frac{\varepsilon_h - 1}{2} + \frac{\varphi_h - 1}{2}}, \quad \left(\frac{\sigma_{h-1} 2^{\omega_h} \sigma_{h+1}}{\varepsilon_h \varphi_h} \right) = \left(\frac{\sigma_{h-1} 2^{\omega_h} \sigma_{h+1}}{\varphi_h} \right),$$

$$\left(\frac{e_h}{\varepsilon_h \varphi_h} \right) = (-1)^{\frac{e_h - 1}{2} \cdot \frac{\varphi_h - 1}{2}} \left(\frac{\varphi_h}{e_h} \right)$$

die Form annimmt

$$(36) \quad \left(\frac{\varphi_{h-1}}{\varepsilon_h \varphi_h} \right) \cdot \left(\frac{\varphi_{h+1}}{\varepsilon_h \varphi_h} \right) = (-1)^{\frac{\varepsilon_h - 1}{2}} \left(\frac{\sigma_{h-1} 2^{\omega_h} \sigma_{h+1}}{\varphi_h} \right) \cdot (-1)^{\frac{e_h + 1}{2} \cdot \frac{\varphi_h - 1}{2}} \left(\frac{\varphi_h}{e_h} \right).$$

Wir gewinnen demnach die $n + 1$ Formeln

$$\left(\frac{\varphi_1}{\varepsilon_0 \varphi_0} \right) = 1,$$

$$\left(\frac{\varphi_0}{\varepsilon_1 \varphi_1} \right) \cdot \left(\frac{\varphi_2}{\varepsilon_1 \varphi_1} \right) = (-1)^{\frac{\varepsilon_1 - 1}{2}} \left(\frac{\sigma_0 2^{\omega_1} \sigma_2}{\varphi_1} \right) (-1)^{\frac{e_1 + 1}{2} \cdot \frac{\varphi_1 - 1}{2}} \left(\frac{\varphi_1}{e_1} \right),$$

$$\left(\frac{\varphi_1}{\varepsilon_2 \varphi_2} \right) \cdot \left(\frac{\varphi_3}{\varepsilon_2 \varphi_2} \right) = (-1)^{\frac{\varepsilon_2 - 1}{2}} \left(\frac{\sigma_1 2^{\omega_2} \sigma_3}{\varphi_2} \right) (-1)^{\frac{e_2 + 1}{2} \cdot \frac{\varphi_2 - 1}{2}} \left(\frac{\varphi_2}{e_2} \right),$$

$$\dots \dots \dots$$

$$\left(\frac{\varphi_{n-3}}{\varepsilon_{n-2} \varphi_{n-2}} \right) \cdot \left(\frac{\varphi_{n-1}}{\varepsilon_{n-2} \varphi_{n-2}} \right) = (-1)^{\frac{\varepsilon_{n-2} - 1}{2}} \left(\frac{\sigma_{n-3} 2^{\omega_{n-2}} \sigma_{n-1}}{\varphi_{n-2}} \right) (-1)^{\frac{e_{n-2} + 1}{2} \cdot \frac{\varphi_{n-2} - 1}{2}} \left(\frac{\varphi_{n-2}}{e_{n-2}} \right),$$

$$\left(\frac{\varphi_{n-2}}{\varepsilon_{n-1} \varphi_{n-1}} \right) \cdot \left(\frac{\varphi_n}{\varepsilon_{n-1} \varphi_{n-1}} \right) = (-1)^{\frac{\varepsilon_{n-1} - 1}{2}} \left(\frac{\sigma_{n-2} 2^{\omega_{n-1}} \sigma_n}{\varphi_{n-1}} \right) (-1)^{\frac{e_{n-1} + 1}{2} \cdot \frac{\varphi_{n-1} - 1}{2}} \left(\frac{\varphi_{n-1}}{e_{n-1}} \right),$$

$$\left(\frac{\varphi_{n-1}}{\varepsilon_n \varphi_n} \right) = 1.$$

Wir multiplizieren jetzt diese $n + 1$ Gleichungen miteinander und benutzen dabei das Reziprozitätsgesetz, welches sich für zwei ungerade, zueinander relativ prime Zahlen U und V , deren Vorzeichen gleich u und v sind, in der Formel

$$\left(\frac{V}{uU}\right) \cdot \left(\frac{U}{vV}\right) = (-1)^{\frac{u-1}{2} \cdot \frac{v-1}{2} + \frac{U-1}{2} \cdot \frac{V-1}{2}}$$

ausspricht. Dadurch erhalten wir die Relation

$$(37) \left\{ \begin{aligned} (-1)^{\left[\frac{I}{2}\right]} &= (-1)^{\sum_{h=1}^{n-1} \frac{e_h-1}{2}} \cdot (-1)^{\frac{e_0-1}{2} \cdot \frac{e_1-1}{2} + \frac{e_1-1}{2} \cdot \frac{e_2-1}{2} + \dots + \frac{e_{n-2}-1}{2} \cdot \frac{e_{n-1}-1}{2} + \frac{e_{n-1}-1}{2} \cdot \frac{e_n-1}{2}} \\ &= \prod_{h=1}^{n-1} \left(\frac{\sigma_{h-1} 2^{\omega_h} \sigma_{h+1}}{\varphi_h} \right) \cdot \prod_{h=1}^{n-1} \left(\frac{\varphi_h}{e_h} \right) \cdot (-1)^{\sum_{h=1}^{n-1} \frac{e_h+1}{2} \cdot \frac{\varphi_h-1}{2}} \\ &\quad \cdot (-1)^{\frac{\varphi_0-1}{2} \cdot \frac{\varphi_1-1}{2} + \frac{\varphi_1-1}{2} \cdot \frac{\varphi_2-1}{2} + \dots + \frac{\varphi_{n-2}-1}{2} \cdot \frac{\varphi_{n-1}-1}{2} + \frac{\varphi_{n-1}-1}{2} \cdot \frac{\varphi_n-1}{2}} \end{aligned} \right.$$

Bilden wir das Produkt aus den ersten η_k oder aus den ersten $\eta_k + 1$ der obigen Gleichungen, so erkennen wir, daß der Charakter $\mathfrak{G}_k(4)$ mit Hilfe der Charaktere $C(p)$, $C(4)$, $C(8)$ und der Einheit

$$D_k^- = \left(\frac{\varphi_{\eta_k-1}}{\varepsilon_{\eta_k} \varphi_{\eta_k}} \right) (-1)^{\frac{I_{\eta_k}(I_{\eta_k}-1)}{2}}$$

oder der Einheit

$$D_k^+ = \left(\frac{\varphi_{\eta_k+1}}{\varepsilon_{\eta_k} \varphi_{\eta_k}} \right) (-1)^{\frac{I_{\eta_k}(I_{\eta_k}+1)}{2}}$$

dargestellt werden kann. Infolgedessen können wir statt der Charaktere $\mathfrak{G}_k$ die Charaktere D_k^- oder D_k^+ einführen.

Zwischen den beiden Charakteren D_k^- und D_k^+ besteht zufolge der Gleichung (36) die Relation

$$(38) \quad D_k^- \cdot D_k^+ = \left(\frac{\sigma_{\eta_k-1} \cdot 2^{\omega_{\eta_k}} \cdot \sigma_{\eta_k+1}}{\varphi_{\eta_k}} \right) (-1)^{\frac{e_{\eta_k}+1}{2} \cdot \frac{\varphi_{\eta_k}-1}{2}} \left(\frac{\varphi_{\eta_k}}{e_{\eta_k}} \right).$$

Wir können jetzt den Satz aussprechen:

Eine charakteristische Form φ der primitiven Formenklasse f besitzt:

I. wenn $\sigma_{h-1} \sigma_h \sigma_{h+1} \equiv 0 \pmod{p}$ ist, einen Charakter

$$\left(\frac{\varphi_h}{p} \right),$$

II. wenn $\sigma_{h-1} \sigma_h \sigma_{h+1} \equiv 0 \pmod{4}$ ist, drei Charaktere

$$(-1)^{\frac{\varphi_h-1}{2}}, \quad \left(\frac{\varphi_{h-1}}{\varepsilon_h \varphi_h} \right) (-1)^{\frac{I_h(I_h-1)}{2}}, \quad \left(\frac{\varphi_{h+1}}{\varepsilon_h \varphi_h} \right) (-1)^{\frac{I_h(I_h+1)}{2}},$$

III. wenn $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{8}$ ist, einen Charakter

$$\left(\frac{2}{\varphi_h}\right).$$

Diese Charaktere sind an die folgenden Bedingungen gebunden:

1. Es sind

$$\left(\frac{\varphi_0}{p}\right) = 1, \quad \left(\frac{\varphi_n}{p}\right) = (-1)^{\frac{p-1}{2} \cdot I}; \quad (-1)^{\frac{\varphi_0-1}{2}} = 1, \quad (-1)^{\frac{\varphi_n-1}{2}} = (-1)^I;$$

$$\left(\frac{2}{\varphi_0}\right) = 1, \quad \left(\frac{2}{\varphi_n}\right) = 1;$$

$$\left(\frac{\varphi_1}{\varepsilon_0 \varphi_0}\right) (-1)^{\frac{I_0(I_0+1)}{2}} = 1, \quad \left(\frac{\varphi_{n-1}}{\varepsilon_n \varphi_n}\right) (-1)^{\frac{I_n(I_n-1)}{2}} = (-1)^{\left[\frac{I}{2}\right]}.$$

2. Wenn $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{4}$ ist, so wird

$$\left[\left(\frac{\varphi_{h-1}}{\varepsilon_h \varphi_h}\right) (-1)^{\frac{I_h(I_h-1)}{2}}\right] \left[\left(\frac{\varphi_{h+1}}{\varepsilon_h \varphi_h}\right) (-1)^{\frac{I_h(I_h+1)}{2}}\right] = \left(\frac{\sigma_{h-1} 2^{\omega_h} \sigma_{h+1}}{\varphi_h}\right) (-1)^{\frac{e_h+1}{2} \cdot \frac{\varphi_h-1}{2}} \left(\frac{\varphi_h}{e_h}\right).$$

3. Wenn $\sigma_h = 2$ ist, in welchem Falle

$$\sigma_{h-2} o_{h-1} \sigma_h \equiv 0 \equiv \sigma_h o_{h+1} \sigma_{h+2} \pmod{4} \quad \text{und} \quad o_h \equiv 1 \pmod{2}$$

wird, so hat zwischen den Charakteren

$$(-1)^{\frac{\varphi_{h-1}-1}{2}} \quad \text{und} \quad (-1)^{\frac{\varphi_{h+1}-1}{2}}$$

die Beziehung statt:

$$(-1)^{\frac{\varphi_{h-1}-1}{2}} \cdot (-1)^{\frac{\varphi_{h+1}-1}{2}} = (-1)^{\frac{e_h+1}{2}}.$$

4. Wenn $\sigma_{h-2} o_{h-1} \sigma_h \equiv 0 \equiv \sigma_h o_{h+1} \sigma_{h+2} \pmod{4}$ und $o_h \equiv 1 \pmod{2}$

ist und zwischen den Charakteren $(-1)^{\frac{\varphi_{h-1}-1}{2}}$ und $(-1)^{\frac{\varphi_{h+1}-1}{2}}$ die Gleichung

$$(-1)^{\frac{\varphi_{h-1}-1}{2}} \cdot (-1)^{\frac{\varphi_{h+1}-1}{2}} = (-1)^{\frac{e_h+1}{2}}$$

statthat, so wird

$$\left[\left(\frac{\varphi_h}{\varepsilon_{h-1} \varphi_{h-1}}\right) (-1)^{\frac{I_{h-1}(I_{h-1}+1)}{2}}\right] \left[\left(\frac{\varphi_h}{\varepsilon_{h+1} \varphi_{h+1}}\right) (-1)^{\frac{I_{h+1}(I_{h+1}-1)}{2}}\right] = \left(\frac{\varphi_h}{e_h}\right).$$

5. Wenn $\sigma_{h,-1} o_{h_i} \sigma_{h_i+1} \equiv 0 \pmod{4}$, $\sigma_{h_i,-1} o_{h_i} \sigma_{h_i+1} \equiv 0 \pmod{4}$ und $h_i - h_{ii} = \pm 1$ ist, so wird

$$\left[\left(\frac{\varphi_{h_i}}{\varepsilon_{h_i} \varphi_{h_i}}\right) (-1)^{\frac{I_{h_i}[I_{h_i}+(h_i-h_{ii})]}{2}}\right] \left[\left(\frac{\varphi_{h_{ii}}}{\varepsilon_{h_i} \varphi_{h_i}}\right) (-1)^{\frac{I_{h_i}[I_{h_i}+(h_{ii}-h_i)]}{2}}\right] = (-1)^{\frac{\varphi_{h_i}-1}{2} \cdot \frac{\varphi_{h_{ii}}-1}{2}}.$$

Die Bedingungen 1. ergeben sich aus den Beziehungen $\varphi_0 = 1$, $\varphi_n = (-1)^I$; die Bedingung 2. stimmt mit der Gleichung (38) überein;

die Bedingung 3. entspringt aus der Kongruenz (35), welche für ein $\sigma_h = 2$ die Relation

$$-e_h \varphi_{h-1} \varphi_{h+1} \equiv 1 \pmod{4}$$

liefert; die Bedingung 4. schließen wir aus der Gleichung

$$\begin{aligned} & \left[\left(\frac{\varphi_h}{\varepsilon_{h-1} \varphi_{h-1}} \right) (-1)^{\frac{I_{h-1}(I_{h-1}-1)}{2}} \right] \left[\left(\frac{\varphi_h}{\varepsilon_{h+1} \varphi_{h+1}} \right) (-1)^{\frac{I_{h+1}(I_{h+1}-1)}{2}} \right] \\ &= \left(\frac{\sigma_{h-1} 2^{\omega_h} \sigma_{h+1}}{\varphi_h} \right) (-1)^{\frac{e_h \varphi_{h-1} \varphi_{h+1} + 1}{2} \cdot \frac{\varphi_h - 1}{2}} \left(\frac{\varphi_h}{e_h} \right), \end{aligned}$$

welche eine unmittelbare Folge von (36) ist; die Bedingung 5. ist eine Identität. — Drücken wir die Charaktere D_k durch die Charaktere $\mathfrak{G}_k(4)$ aus, so werden sämtliche Bedingungen 2., 4., 5. zu Identitäten, und es bleiben allein die Bedingungen 3. und die Bedingung (37) zu erfüllen.

II. Zu jeder Kombination von Charakteren I., II., III., welche allen Bedingungen 1., 2., 3., 4., 5. genügt, gehört ein Genus der Ordnung O , welches wirklich primitive Formen enthält.

Beweis: Es möge irgendeine Kombination der Charaktere I., II., III. oder der Charaktere $C(p)$, $C(4)$, $C(8)$, $\mathfrak{G}(4)$ gegeben sein, welche keiner der aufgestellten Bedingungen widerspricht. Alsdann können wir leicht $n-1$ Zahlen $\bar{\varphi}_h$ finden, welche zu den Zahlen $2o_h$ relativ prim sind und für welche die Einheiten $C_p(\bar{\varphi})$, $C_4(\bar{\varphi})$, $C_8(\bar{\varphi})$, $\mathfrak{G}_4(\bar{\varphi})$ gleich den gegebenen Größen $C(p)$, $C(4)$, $C(8)$, $\mathfrak{G}(4)$ ausfallen. Da die Einheiten $C_p(\bar{\varphi})$, $C_4(\bar{\varphi})$, $C_8(\bar{\varphi})$, $\mathfrak{G}_4(\bar{\varphi})$ nur von den Resten der Zahlen $\bar{\varphi}_h$ in bezug auf die Moduln $8o_h$ abhängen, so erhellt, daß für irgendein System von Zahlen φ_h , welches den Kongruenzen

$$\varphi_h \equiv \bar{\varphi}_h \pmod{8o_h} \quad (h = 1, 2, \dots, n-2, n-1)$$

genügt, gewiß auch die Einheiten $C_p(\varphi)$, $C_4(\varphi)$, $C_8(\varphi)$, $\mathfrak{G}_4(\varphi)$ gleich den gegebenen Einheiten $C(p)$, $C(4)$, $C(8)$, $\mathfrak{G}(4)$ sein werden.

Wir wählen nun n Einheiten $\delta_1, \delta_2, \dots, \delta_n$, von welchen I gleich -1 und $n-I$ gleich $+1$ sein mögen, und setzen $\delta_1 \delta_2 \dots \delta_n = \varepsilon_h$. Aus dem bekannten Satz, daß eine jede arithmetische Progression $zS + s$ (s relativ prim zu S , $z = -\infty, \dots, -2, -1, 0, 1, 2, \dots, +\infty$) sowohl unendlich viele positive als unendlich viele negative Primzahlen enthält, erkennen wir (mit Hinzuziehung eines einfachen Schlusses von $h-1$ auf h), daß wir die Zahlen $\varphi_h \equiv \bar{\varphi}_h \pmod{8o_h}$ ($h = 1, 2, \dots, n-2, n-1$) derart bestimmen können, daß die Größen $\varepsilon_h \varphi_h$ positive Primzahlen werden, daß die Zahlen φ_h zu den Zahlen $2o_1 o_2 \dots o_{n-2} o_{n-1} \cdot \varphi_{h-1}$ relativ prim ausfallen und daß $n-1$ Kongruenzen der Form

$$[\varphi_0 = 1] \quad -\sigma_{h-1} o_h \sigma_{h+1} \varphi_{h-1} \varphi_{h+1} \equiv X_h^2 \pmod{\sigma_h^2 \varphi_h} \quad [\varphi_n = (-1)^I]$$

statthaben.*) Nachdem wir $n - 1$ derartige Zahlen φ_h gefunden haben, suchen wir $n - 1$ Zahlen $y_1, y_2, \dots, y_{n-1}$ auf, welche den Kongruenzen

$$[y_0 = 0] \quad -\frac{1}{o_h} \frac{\sigma_{h+1} \varphi_{h+1}}{\sigma_{h-1} \varphi_{h-1}} \equiv y_h^2 \pmod{\sigma_h \varphi_h} \quad (h = 1, 2, \dots, n-1)$$

genügen, und wir setzen die ganzen Zahlen

$$\frac{\sigma_{h+1} \varphi_{h+1} + \sigma_{h-1} \varphi_{h-1} o_h y_h^2}{\sigma_h \varphi_h} = t_h$$

und

$$\text{Die Form} \quad o_1 o_2 \cdots o_h \cdot y_h = Y_h, \quad o_1 o_2 \cdots o_h \cdot t_h = T_h.$$

$$\varphi = \begin{pmatrix} T_0, & Y_1 \\ Y_1, & T_1, & Y_2 \\ & Y_2, & T_2, & Y_3 \\ & & \ddots & \ddots & \ddots \\ & & & Y_{n-2}, & T_{n-2}, & Y_{n-1} \\ & & & & Y_{n-1}, & T_{n-1} \end{pmatrix}$$

bildet jetzt, wie man sich leicht überzeugt, eine charakteristische Form für ein Genus G , welches die gegebenen Charaktere besitzt. — Die aus den ersten h Reihen von φ gebildeten symmetrischen Unterdeterminanten werden gleich den Zahlen $\sigma_h d_{h-1} \varphi_h$.

Das soeben entwickelte Verfahren dient, wie zum Beweise der Existenz eines Formengenus G , so auch zum Nachweise der Existenz einer beliebigen Ordnung. Um ein Beispiel zu geben, zeigen wir, daß es primitive Formen mit 8 Variablen von der Determinante 1 gibt, welche der Ordnung

$$\begin{pmatrix} \sigma_h \\ o_h \end{pmatrix} = \begin{pmatrix} 2, 1, 2, 1, 2, 1, 2 \\ 1, 1, 1, 1, 1, 1, 1 \end{pmatrix}, \quad I = 0$$

angehören.

Man erkennt, daß die $8 - 1 = 7$ Zahlen

$$\varphi_1 = 1, \varphi_2 = 3, \varphi_3 = 5, \varphi_4 = 13, \varphi_5 = 5, \varphi_6 = 3, \varphi_7 = 1$$

den Kongruenzen (36) genügen, und man erhält die Form

$$\varphi^{(8)} = \begin{pmatrix} 2, & 1 \\ 1, & 2, & 1 \\ & 1, & 4, & 3 \\ & & 3, & 4, & 5 \\ & & & 5, & 20, & 3 \\ & & & & 3, & 12, & 1 \\ & & & & & 1, & 4, & 1 \\ & & & & & & 1, & 2 \end{pmatrix} \sim \begin{pmatrix} 2, & 1 \\ 1, & 2, & 1, & 0, & -1 \\ & 1, & 2, & 1, & 0 \\ & & 0, & 1, & 2, & 1 \\ & & & -1, & 0, & 1, & 2, & 1 \\ & & & & & 1, & 2, & 1 \\ & & & & & & 1, & 2, & 1 \\ & & & & & & & 1, & 2 \end{pmatrix}$$

*) Siehe Dirichlet, Vorlesungen über Zahlentheorie, herausgeg. von Dedekind, 1880, S. 328. [[4. Auflage (1894), S. 328]].

Da diese Form einer anderen Ordnung angehört wie die Form $\varphi^{(0)} = \sum_{h=1}^8 x_h^2$, so kann sie dieser Form $\varphi^{(0)}$ nicht äquivalent sein, und wir sehen somit, daß die Formen von 8 Variablen und der Determinante 1 mindestens zwei Formenklassen liefern. Entsprechend liefern die Formen von $n \left(= 8 \left[\frac{n}{8} \right] + n_0, n_0 < 8 \right)$ Variablen und der Determinante 1 mindestens $\left[\frac{n}{8} \right] + 1$ Formenklassen. Denn es sind die $\left[\frac{n}{8} \right] + 1$ Formen

$$\varphi_{(m)} = \varphi^{(8)}(x_1, \dots, x_8) + \varphi^{(8)}(x_9, \dots, x_{16}) + \dots + \varphi^{(8)}(x_{8m-7}, \dots, x_{8m}) + \sum_{h=8m+1}^n x_h^2$$

$$(m = 0, 1, 2, \dots, \left[\frac{n}{8} \right])$$

sämtlich von der Determinante 1, und es können nicht zwei von diesen Formen einander äquivalent sein, da die Anzahl der Darstellungen der Zahl 1 durch eine dieser Formen $\varphi_{(m)}$ gleich $2(n - 8m)$ ist und mithin für keine zwei dieser Formen denselben Wert erhält.

III. Die Anzahl der sämtlichen Kombinationen der Charaktere I., II., III., welche mit den aufgestellten Bedingungen verträglich sind, möge durch g bezeichnet werden. Wir wollen die Anzahl der Größen $\sigma_{h-1} o_h \sigma_{h+1}$ ($h = 1, 2, \dots, n-2, n-1$), welche durch eine ungerade Primzahl p oder durch die Zahl 4 oder durch die Zahl 8 teilbar sind, von neuem gleich $\lambda_p - 1$, gleich $\mu - 1$, gleich $\nu - 1$ setzen; die Anzahl der Fälle, in welchen zwei aufeinanderfolgende Zahlen $\sigma_{h-1} o_h \sigma_{h+1}$ ($h = 0, 1, 2, \dots, n-1, n$) gleichzeitig durch 4 teilbar sind, sei gleich μ_1 , und die Anzahl der Fälle, in welchen $\sigma_h = 1$, $\sigma_{h-2} o_{h-1} \sigma_h \equiv 0 \equiv \sigma_h o_{h+1} \sigma_{h+2} \pmod{4}$, $o_h \equiv 1 \pmod{2}$ ($h = 1, 2, \dots, n-1$) ist, sei gleich μ_2 ; ferner möge $\mu - \mu_0$ die Anzahl aller Invarianten σ_h ($h = 1, 2, \dots, n-1$) bedeuten, welche gleich 2 sind.

Die Zahl g , welche zugleich die Anzahl der sämtlichen wirklich vorhandenen Genera von der Ordnung

$$O: \begin{pmatrix} \sigma_1, \sigma_2, \dots, \sigma_{n-2}, \sigma_{n-1} \\ o_1, o_2, \dots, o_{n-2}, o_{n-1} \end{pmatrix}, \quad I$$

ergibt, erhält den Ausdruck

$$g = g_0 [1 + \{\mu_0 - \mu_2\} + \square(\mu_0 - \mu_1 - \mu_2)],$$

in welchem die Größen g_0 , $\{\mu_0 - \mu_2\}$, $\square(\mu_0 - \mu_1 - \mu_2)$ die nachstehende Bedeutung haben:

Es ist

$$g_0 = 2^{\sum (\lambda_p - 1) + 2(\mu_0 - 1) + (\nu - 1)} \left(\frac{1}{2} \right)^{\mu_1} \left(\frac{3}{4} \right)^{\mu_2},$$

worin die Summe $\sum_{(p)} (\lambda_p - 1)$ über alle Größen λ_p zu erstrecken ist, welche einer in dem Produkt $\prod_{h=1}^{n-1} o_h$ aufgehenden ungeraden Primzahl p entsprechen.

Es ist

$$\{\mu_0 - \mu_n\} = 0,$$

sobald $\mu_0 - \mu_n > 0$ ist, und

$$\{\mu_0 - \mu_n\} = (-1)^{\frac{n}{2} - I + \mu_n} \cdot \left(\frac{-1}{\prod_{h=1}^{\frac{n}{2}} \frac{o_{2h-1}}{2^{\omega_{2h-1}}}} \right) \cdot \frac{1}{3^{\mu_n}}, \quad [n \equiv 0 \pmod{2}]$$

sobald $\mu_0 - \mu_n = 0$, d. h. $\mu_n = 0$, $n = 2\mu$ ist.

Es wird

$$\square (\mu_0 - \mu_n - \mu_n) = 0,$$

wenn $\mu_0 - \mu_n - \mu_n > 0$ ist oder wenn irgendeine der Zahlen $\sigma_{h-1} o_h \sigma_{h+1}$ ($h = 1, 2, \dots, n-1$) kein Quadrat ist, dagegen wird

$$\square (\mu_0 - \mu_n - \mu_n) = \delta(n - 2I) \cdot \frac{1}{3^{\mu_n} 2^{\left\lfloor \frac{\mu_n - 1}{2} \right\rfloor}}$$

$$\left\{ \begin{array}{l} n - 2I \equiv 1 \mid 3 \mid 5 \mid 7 \mid 0 \mid 2 \mid 4 \mid 6 \pmod{8} \\ \delta(n - 2I) = 1 \mid -1 \mid -1 \mid 1 \mid 1 \mid 0 \mid -1 \mid 0 \end{array} \right\},$$

wenn $\mu_0 - \mu_n - \mu_n = 0$ ist und die sämtlichen Zahlen $\sigma_{h-1} o_h \sigma_{h+1}$ Quadrate sind.

Offenbar besitzt nach dem Satze II. eine Ordnung O stets primitive Formen, wenn nicht $g = 0$ ist. Man erkennt leicht, daß nur in den folgenden Fällen $g = 0$ werden kann:

($\square = 0$) wenn

$$\mu_0 = 0; \mu_n = 0, \mu_n = 0, \quad (-1)^{\frac{n}{2} - I} \prod_{h=1}^{\frac{n}{2}} \frac{o_{2h-1}}{2^{\omega_{2h-1}}} \equiv -1 \pmod{4}$$

ist;

($\delta = -1$) wenn alle Zahlen $\sigma_{h-1} o_h \sigma_{h+1}$ Quadratzahlen sind und entweder

$$\mu_0 = 0; \mu_n = 0, \mu_n = 0, \quad n - 2I \equiv 4 \pmod{8}$$

oder

$$\mu_0 = 1; \mu_n = 1, \mu_n = 0, \quad n - 2I \equiv 3, 5 \pmod{8}$$

oder

$$\mu_0 = 1; \mu_n = 0, \mu_n = 1, \quad n - 2I \equiv 4 \pmod{8}$$

oder

$$\mu_0 = 2; \mu_n = 2, \mu_n = 0, \quad n - 2I \equiv 4 \pmod{8}$$

ist.

Kap. XII. Adjungierte Formen. — Reziprozität zwischen den

$$\text{Ordnungen } \begin{pmatrix} \sigma_1, & \sigma_2, & \dots, & \sigma_{n-2}, & \sigma_{n-1} \\ \rho_1, & \rho_2, & \dots, & \rho_{n-2}, & \rho_{n-1} \end{pmatrix}, I$$

$$\text{und } \begin{pmatrix} \sigma_{n-1}, & \sigma_{n-2}, & \dots, & \sigma_2, & \sigma_1 \\ \rho_{n-1}, & \rho_{n-2}, & \dots, & \rho_2, & \rho_1 \end{pmatrix}, I.$$

Es sei $f = \sum_{i,k=1}^n a_{ik} x_i x_k$ eine primitive quadratische Form von der Determinante $\Delta(f)$. Der größte gemeinsame Teiler aller $(n-1)$ -reihigen Unterdeterminanten $\frac{\partial \Delta(f)}{\partial a_{ik}}$ ist gleich d_{n-2} . Setzen wir also

$$\frac{1}{d_{n-2}} \cdot \frac{\partial \Delta(f)}{\partial a_{ik}} = \varepsilon \cdot a'_{n-i+1, n-k+1},$$

wo $\varepsilon = (-1)^{I(f)}$ ist, so wird die Form

$$f' = \sum_{i,k=1}^n a'_{ik} x'_i x'_k$$

ebenso wie die Form f primitiv sein. Diese Form f' soll der Form f *adjungiert* heißen, und wir schreiben $f \asymp f'$.

N. Wenn die Form f einer Form $g = \sum_{i,m=1}^n b_{im} y_i y_m$ äquivalent ist, so ist die zu f adjungierte Form f' der zu g adjungierten Form $g' = \sum_{i,m=1}^n b'_{im} y'_i y'_m$ äquivalent.

Denn nehmen wir an, die Form f gehe in g durch eine Substitution

$$S: \quad x_i = \sum_{l=1}^n s_i^l y_l$$

über, so wird

$$B \begin{pmatrix} l_1, & l_2, & \dots, & l_{n-1} \\ m_1, & m_2, & \dots, & m_{n-1} \end{pmatrix} \\ = \sum A \begin{pmatrix} i_1, & i_2, & \dots, & i_{n-1} \\ k_1, & k_2, & \dots, & k_{n-1} \end{pmatrix} S \begin{pmatrix} l_1, & l_2, & \dots, & l_{n-1} \\ i_1, & i_2, & \dots, & i_{n-1} \end{pmatrix} S \begin{pmatrix} m_1, & m_2, & \dots, & m_{n-1} \\ k_1, & k_2, & \dots, & k_{n-1} \end{pmatrix}.$$

Die Unterdeterminanten $S \begin{pmatrix} l_1, & l_2, & \dots, & l_{n-1} \\ i_1, & i_2, & \dots, & i_{n-1} \end{pmatrix}$ stimmen dem absoluten Werte nach mit den Zahlen $\frac{\partial |S|}{\partial s_i^l} = s'^{n-l+1}_{n-i+1}$ überein, und wir erkennen leicht, daß die vorstehende Gleichung sich schreiben läßt

$$(-1)^l d_{n-2} b'_{lm} = \sum_{i,k=1}^n (-1)^l d_{n-2} a'_{ik} s_i^l s_k'^m.$$

Demnach wird die Form f' in g' durch die Substitution

$$S': \quad x'_i = \sum_{l=1}^n s'_i{}^l y'_l$$

transformiert. Man findet noch $|S'| = |S|^{n-1}$. Wenn also $|S| = 1$ und $f \sim g$ ist, so wird $|S'| = 1$ und $f' \sim g'$.

Die Substitution S' möge der Substitution S adjungiert heißen ($S \times S'$).

Wenn f vermittelt einer linearen Substitution in eine Summe von n Quadraten

$$f = - \sum_{h=1}^{I(f)} x_h^2 + \sum_{h=1}^{n-I(f)} \Xi_h^2$$

transformiert wird, so läßt sich die Form f' in der Gestalt

$$f' = - \sum_{h=1}^{I(f)} x'_h{}^2 + \sum_{h=1}^{n-I(f)} \Xi'_h{}^2$$

schreiben. Es gilt also die Beziehung

$$I(f) = I(f') = I.$$

Wir bezeichnen die $n-1$ Invarianten σ, o, d der Form f' durch σ'_h, o'_h, d'_h ($h = 1, 2, \dots, n-2, n-1$).

Man beweist leicht den folgenden Satz:

Wenn f' der Form f adjungiert ist, wird auch f der Form f' adjungiert sein.

Denn ist $f'' = \sum_{i,k=1}^n a''_{ik} x''_i x''_k$ die zu f' adjungierte Form, so gelten nach einem bekannten Determinantensatz die Gleichungen

$$(\varepsilon d_{n-2})^{n-1} \cdot \varepsilon d'_{n-2} a''_{ik} = (\varepsilon d_{n-2})^{n-1} \cdot \frac{\partial \Delta(f')}{\partial a'_{n-i+1, n-k+1}} = (\varepsilon d_{n-1})^{n-2} \cdot a_{ik},$$

in denen $\varepsilon = (-1)^I$ ist. Demnach ist der Quotient $\frac{a''_{ik}}{a_{ik}}$ immer positiv und für alle Werte von i und k derselbe. Da aber die Größen a''_{ik} ebenso wie die a_{ik} ganze Zahlen ohne einen gemeinsamen Teiler > 1 sind, so muß $\frac{a''_{ik}}{a_{ik}} = 1$, d. i. $f'' = f$ werden.

Man kann die Ordnung der Form f' aus der Ordnung der Form f herleiten. Wenn wir die symmetrischen h -reihigen Unterdeterminanten der Form f und der Form f' durch F_h und durch F'_h bezeichnen, die unsymmetrischen aber durch P_h und durch P'_h , so bestehen nach einem bekannten Satz zwischen den F_h, P_h, F'_h, P'_h die Relationen

$$(\varepsilon d_{n-2})^h F'_h = (\varepsilon d_{n-1})^{h-1} F_{n-h}, \quad (\varepsilon d_{n-2})^h P'_h = (\varepsilon d_{n-1})^{h-1} P_{n-h}.$$

Infolgedessen muß erstens der größte positive Teiler aller Zahlen $(\varepsilon d_{n-2})^h F'_h$,

$(\varepsilon d_{n-2})^h P'_h$ mit dem größten positiven Teiler der Zahlen $(\varepsilon d_{n-1})^{h-1} F_{n-h}$, $(\varepsilon d_{n-1})^{h-1} P_{n-h}$ übereinstimmen; zweitens wird der größte positive Teiler der Zahlen $(\varepsilon d_{n-2})^h F'_h$, $(\varepsilon d_{n-2})^h 2 P'_h$ dem größten positiven Teiler der Zahlen $(\varepsilon d_{n-1})^{h-1} F_{n-h}$, $(\varepsilon d_{n-1})^{h-1} 2 P_{n-h}$ gleich sein müssen. Wir bekommen demnach die Gleichungen

$$\begin{aligned} d_{n-2}^h d'_{h-1} &= d_{n-1}^{h-1} d_{n-h-1}, \\ \sigma'_h d_{n-2}^h d'_{h-1} &= \sigma_{n-h} d_{n-1}^{h-1} d_{n-h-1}. \end{aligned}$$

Die Division der zweiten Gleichung durch die erste gibt zunächst

$$\sigma'_h = \sigma_{n-h},$$

d. i.

$$(39) \quad \sigma'_1 = \sigma_{n-1}, \sigma'_2 = \sigma_{n-2}, \dots, \sigma'_{n-2} = \sigma_2, \sigma'_{n-1} = \sigma_1.$$

Dann führt die Kombination der drei Gleichungen

$$(40_1) \quad d_{n-2}^{h+1} d'_h = d_{n-1}^h d_{n-h-2},$$

$$(40_2) \quad d_{n-2}^h d'_{h-1} = d_{n-1}^{h-1} d_{n-h-1},$$

$$(40_3) \quad d_{n-2}^{h-1} d'_{h-2} = d_{n-1}^{h-2} d_{n-h},$$

zu

$$\frac{d'_h \cdot d'_{h-2}}{(d'_{h-1})^2} = \frac{d_{n-h} \cdot d_{n-h-2}}{d_{n-h-1}^2}, \quad o'_h = o_{n-h},$$

d. i.

$$(41) \quad o'_1 = o_{n-1}, o'_2 = o_{n-2}, \dots, o'_{n-2} = o_2, o'_{n-1} = o_1.$$

Die Form f' gehört daher einer Ordnung

$$O': \begin{pmatrix} \sigma_{n-1}, \sigma_{n-2}, \dots, \sigma_2, \sigma_1 \\ o_{n-1}, o_{n-2}, \dots, o_2, o_1 \end{pmatrix}, \quad I$$

an.

Unter Benutzung des Satzes N. können wir jetzt das Resultat aussprechen:

Jeder Klasse f der Ordnung $\begin{pmatrix} \sigma_h \\ o_h \end{pmatrix}$, I ist eine bestimmte Klasse f' der

Ordnung $\begin{pmatrix} \sigma_{n-h} \\ o_{n-h} \end{pmatrix}$, I adjungiert.

Wenn φ eine Grundform der Klasse f für den Modul N bedeutet, so stellt die der Form φ adjungierte Form φ' eine Grundform der Klasse f' für den Modul N vor. Denn offenbar sind die $n+1$ Zahlen φ'_h der Form φ' mit den $n+1$ Zahlen φ_h der Form φ durch die Gleichungen

$$\varphi'_h = (-1)^I \varphi_{n-h},$$

d. i.

$\varphi'_1 = (-1)^I \varphi_{n-1}, \varphi'_2 = (-1)^I \varphi_{n-2}, \dots, \varphi'_{n-2} = (-1)^I \varphi_2, \varphi'_{n-1} = (-1)^I \varphi_1$ verbunden.

Demnach sind die Charaktere der Klasse f' aus denen der Klasse f ableitbar. Wir schließen hieraus den Satz:

Gehören die beiden Formen f und g einem und demselben Genus an, so sind auch die ihnen adjungierten Formen f' und g' in einem und demselben Genus enthalten.

Infolgedessen werden zwei *Genera* allemal dann *adjungiert* heißen, wenn eine Form des einen Genus einer Form des andern adjungiert ist.

Zweiter Teil.

Über die Darstellung ganzer Zahlen durch quadratische Formen.

Kap. XIII. Hilfssatz.

Es sei ein System von $n \cdot \nu$ ganzen Zahlen u_k^h ($0 \leq k < n$, $0 \leq h < \nu$)

$$(u) = \begin{pmatrix} u_0^0 & u_1^0 & \dots & u_{n-1}^0 \\ u_0^1 & u_1^1 & \dots & u_{n-1}^1 \\ \dots & \dots & \dots & \dots \\ u_0^{\nu-1} & u_1^{\nu-1} & \dots & u_{n-1}^{\nu-1} \end{pmatrix} \quad (\nu < n)$$

gegeben. Falls die sämtlichen ν -reihigen Unterdeterminanten

$$u \begin{pmatrix} 0, 1, \dots, \nu-1 \\ k_0, k_1, \dots, k_{\nu-1} \end{pmatrix} \quad (k = 0, 1, \dots, n-1),$$

welche sich aus diesem Systeme bilden lassen, keinen gemeinsamen Teiler > 1 besitzen, so können wir $n(n-\nu)$ ganze Zahlen u_k^h ($0 \leq k < n$, $\nu \leq h < n$) so finden, daß die n -reihige Determinante

$$|u_k^h| \quad (k, h = 0, 1, \dots, n-1)$$

den Wert 1 erhält.*)

Beweis. — Dieser Satz ist evident, wenn $\nu = 0$ ist; denn in diesem Fall kann man $u_k^h = 1$ ($k = h$), $u_k^h = 0$ ($k \neq h$) nehmen.

Ist $\nu > 0$, so wird der größte gemeinsame Teiler der n Zahlen $u_0^0, u_1^0, \dots, u_{n-1}^0$ in dem größten gemeinsamen Teiler der Determinanten

$$u \begin{pmatrix} 0, 1, \dots, \nu-1 \\ k_0, k_1, \dots, k_{\nu-1} \end{pmatrix}$$

enthalten und folglich gleich 1 sein.

1. Wir beweisen zunächst, daß, wenn der größte gemeinsame Teiler der n ganzen Zahlen $u_0^0, u_1^0, \dots, u_{n-1}^0$ gleich der Einheit ist, stets eine Substitution

*) Gauß, Disquisitiones arithmeticae, art. 279.

$$S: \quad v_i = \sum_k s_i^k u_k \quad (i, k = 0, 1, \dots, n-1)$$

von der Determinante 1 gefunden werden kann, welche den n Gleichungen

$$\sum_k s_0^k u_k^0 = 1, \quad \sum_k s_1^k u_k^0 = 0, \dots, \quad \sum_k s_{n-1}^k u_k^0 = 0$$

genügt, d. h. welche die n Zahlen $u_k = u_k^0$ durch die n Größen $v_0 = 1, v_1 = 0, \dots, v_{n-1} = 0$ ersetzt.

In der Tat, falls von den n Zahlen $u_0^0, u_1^0, \dots, u_{n-1}^0$ nur eine einzige, etwa u_i^0 , von Null verschieden ist, so stellt diese Zahl offenbar zugleich den größten Teiler der sämtlichen n Zahlen u_k^0 vor, und es wird demnach $u_i^0 = \pm 1, \sum_{h=0}^{n-1} (u_h^0)^2 = 1$ sein. Ist dann $i = 0$, so kann man als Substitution S die folgende wählen:

$$v_0 = \pm u_0, \quad v_i = \pm u_i; \quad v_h = u_h \quad (h \neq 0, i_0),$$

in der i_0 irgend einen von 0 verschiedenen Index bedeutet; wenn aber $i > 0$ ist, statt dessen die folgende:

$$v_0 = \pm u_i, \quad v_i = \pm u_0; \quad v_h = u_h \quad (h \neq 0, i).$$

Falls aber unter den n Zahlen $u_0^0, u_1^0, \dots, u_{n-1}^0$ mindestens zwei, etwa u_i^0, u_k^0 von Null verschieden sind, so wird $\sum_{h=0}^{n-1} (u_h^0)^2 > 1$, und wenn $(u_i^0)^2 \geq (u_k^0)^2$ ist, können wir eine Einheit ± 1 so bestimmen, daß $(u_i^0)^2 > (u_i^0 \pm u_k^0)^2$ ausfällt. Durch Ausübung der Substitution

$$(s): \quad U_i^0 = u_i^0 \pm u_k^0, \quad U_k^0 = u_k^0; \quad U_h^0 = u_h^0 \quad (h \neq i, k)$$

von der Determinante 1 erhalten wir dann n Zahlen $U_0^0, U_1^0, \dots, U_{n-1}^0$ ohne gemeinsamen Teiler, für welche die Ungleichung

$$\sum_{h=0}^{n-1} (u_h^0)^2 > \sum_{h=0}^{n-1} (U_h^0)^2$$

statthat.

Nehmen wir an, daß der Punkt 1. unseres Satzes bereits für alle Systeme $U_0^0, U_1^0, \dots, U_{n-1}^0$ ohne gemeinsamen Teiler, für welche die Größe $\sum_{h=0}^{n-1} (U_h^0)^2$ kleiner als $\sum_{h=0}^{n-1} (u_h^0)^2$ ist, bewiesen sei, so können wir eine Substitution (τ) von der Determinante 1 finden, welche die Zahlen $U_0^0, U_1^0, \dots, U_{n-1}^0$ durch die Zahlen $1, 0, \dots, 0$ ersetzt, und die zusammengesetzte Substitution $S = (\tau) \cdot (s)$ ergibt alsdann das gleiche Resultat für die Zahlen $u_0^0, u_1^0, \dots, u_{n-1}^0$.

2. Bilden wir das Produkt des Systemes

$$u = \begin{pmatrix} u_0^0, & u_1^0, & \dots, & u_{n-1}^0 \\ u_0^1, & u_1^1, & \dots, & u_{n-1}^1 \\ \cdot & \cdot & \cdot & \cdot \\ u_0^{n-1}, & u_1^{n-1}, & \dots, & u_{n-1}^{n-1} \end{pmatrix},$$

in welchem die Zahlen u_k^h ($\nu \leq h < n$) vorläufig unbestimmte Größen sein mögen, und des Systemes

$$S = \begin{pmatrix} s_0^0, & s_1^0, & \dots, & s_{n-1}^0 \\ s_0^1, & s_1^1, & \dots, & s_{n-1}^1 \\ \cdot & \cdot & \cdot & \cdot \\ s_0^{n-1}, & s_1^{n-1}, & \dots, & s_{n-1}^{n-1} \end{pmatrix},$$

so ergibt sich ein System von der Gestalt

$$v = u \cdot S = \begin{pmatrix} 1, & 0, & \dots, & 0 \\ v_0^1, & v_1^1, & \dots, & v_{n-1}^1 \\ \cdot & \cdot & \cdot & \cdot \\ v_0^{n-1}, & v_1^{n-1}, & \dots, & v_{n-1}^{n-1} \end{pmatrix},$$

in welchem die Größen v_k^h ($1 \leq h < \nu$) allein von den s_i^h und den gegebenen Zahlen u_k^h abhängen, während die Größen v_k^h ($\nu \leq h < n$) sich durch die s_i^h und die gesuchten Zahlen u_k^h ausdrücken.

Aus dem Umstande, daß die Determinante $|s_i^k|$ gleich 1 ist, erkennen wir, daß der größte gemeinsame Teiler der sämtlichen $(\nu - 1)$ -reihigen Unterdeterminanten

$$v \begin{pmatrix} 1, & \dots, & \nu - 1 \\ (k_1, & \dots, & k_{\nu-1}) \end{pmatrix} \quad (k = 1, \dots, n - 1)$$

des Systems

$$(v) = \begin{pmatrix} v_1^1, & \dots, & v_{n-1}^1 \\ \cdot & \cdot & \cdot \\ v_1^{\nu-1}, & \dots, & v_{n-1}^{\nu-1} \end{pmatrix}$$

dem größten gemeinsamen Teiler der sämtlichen Unterdeterminanten

$$u \begin{pmatrix} 0, & 1, & \dots, & \nu - 1 \\ (k_0, & k_1, & \dots, & k_{\nu-1}) \end{pmatrix},$$

d. i. der Einheit gleich sein wird. Nehmen wir also unsern Satz für den Fall $n - 1$, $\nu - 1$ schon als bewiesen an, so können wir solche ganze Zahlen $v_1^h, \dots, v_{n-1}^h$ ($\nu \leq h < n$) bestimmen, daß die Determinante

$$|v_k^h| \quad (k, h = 1, \dots, n - 1)$$

den Wert 1 erhält. Setzen wir alsdann noch für die Zahlen v_0^h ($\nu \leq h < n$) beliebige ganze Zahlen ein, so besitzt das zusammengesetzte System $v \cdot S^{-1}$

offenbar die Determinante 1 und lauter ganzzahlige Koeffizienten, und seine ν ersten Reihen stimmen mit den ν Reihen des Systems (u) überein. Demnach liefert uns dieses System $\nu \cdot S^{-1} = u$ sofort $n(n - \nu)$ ganze Zahlen $u_k^h (\nu \leq h < n)$, für welche

$$|u_k^h| = 1 \quad (k, h = 0, 1, \dots, n-1)$$

wird, und hierdurch ist unser Satz für den Fall n, ν bewiesen. Da er gewiß für den Fall $n - \nu, 0$ statthat, ergibt sich also seine Gültigkeit für die Fälle

$$n - \nu + 1, 1; n - \nu + 2, 2; \dots; n, \nu.$$

Kap. XIV. Darstellung einer Form von ν Variablen durch eine Form von n Variablen. — Äquivalente Darstellungen und Darstellungsgruppen.

I. Wir sagen, eine Form

$$\varphi = \sum_{i,k=1}^{\nu} \alpha_{ik} \xi_i \xi_k$$

von ν Variablen sei *darstellbar* durch eine Form

$$f = \sum_{i,k=1}^n a_{ik} x_i x_k$$

von n Variablen ($n > \nu$), wenn die Form f vermittle einer Substitution

$$(r): \quad x_i = \sum_{k=1}^{\nu} r_i^k \xi_k \quad (i = 1, 2, \dots, n),$$

in welcher die Größen r_i^k ganze Zahlen sind, in φ übergeht.

Für die Koeffizienten α_{im} ergeben sich die Gleichungen

$$\alpha_{im} = \sum_{i,k=1}^n a_{ik} r_i^i r_k^m.$$

Aus denselben erhellt, daß der größte Teiler d_0 der sämtlichen Koeffizienten a_{ik} in dem größten Teiler der sämtlichen Koeffizienten α_{im} enthalten ist und daß die Substitution (r) , welche die Form φ durch f darstellt, auch zur Darstellung der Form $\frac{\varphi}{d_0} = \left\{ \frac{\alpha_{ik}}{d_0} \right\}$ durch die primitive Form $\frac{f}{d_0} = \left\{ \frac{a_{ik}}{d_0} \right\}$ dienen kann. Infolge dieses Umstandes dürfen wir uns auf die Untersuchung der Fälle, in denen die Form f primitiv ist, beschränken.

Im Folgenden betrachten wir insbesondere Darstellungen (r) , für welche der größte gemeinsame Teiler der ν -reihigen Unterdeterminanten

des Systems

$$\begin{pmatrix} r_1^1 & r_2^1 & \dots & r_n^1 \\ r_1^2 & r_2^2 & \dots & r_n^2 \\ \cdot & \cdot & \cdot & \cdot \\ r_1^\nu & r_2^\nu & \dots & r_n^\nu \end{pmatrix} \quad (\nu < n)$$

gleich 1 ist, und welche wir *eigentliche Darstellungen* nennen.

Wenn die Darstellung (r) eine eigentliche ist, so können wir nach unserm Hilfssatz $n(n - \nu)$ Zahlen r_i^k ($k > \nu$) finden derart, daß die Determinante

$$|r_i^k| \quad (i, k = 1, 2, \dots, n)$$

den Wert 1 erhält. Die Form f geht alsdann durch die Substitution

$$R: x_i = \sum_{k=1}^n r_i^k \xi_k \quad (i = 1, 2, \dots, n)$$

in eine äquivalente Form Φ über, in welcher die aus den ersten ν Horizontal- und Vertikalreihen gebildete Form mit φ identisch ist.

Wenn umgekehrt in der Klasse f ein Repräsentant Φ vorhanden ist, in welchem die aus den ersten ν Reihen gebildete Form gleich φ wird, so ergibt jede Substitution, vermittelt deren f in Φ übergeht, eine eigentliche Darstellung von φ durch f .

Wir gewinnen aus dieser Bemerkung, in welcher die Form f nur als Repräsentant ihrer Klasse erscheint, den Satz:

Durch zwei äquivalente Formen f und g können ebendieselben Formen φ eigentlich dargestellt werden.

Ferner ist es leicht den folgenden Satz zu beweisen:

O. Wenn die Form $\varphi = \sum_{i,k=1}^{\nu} \alpha_{ik} \xi_i \xi_k$ durch die Form f eigentlich dargestellt werden kann, so ist auch jede der Form φ äquivalente Form

$\psi = \sum_{i,k=1}^{\nu} \beta_{ik} \eta_i \eta_k$ durch f eigentlich darstellbar.

In der Tat: es möge φ durch die Substitution

$$(\tau): \xi_i = \sum_{k=1}^{\nu} \tau_i^k \eta_k \quad (i = 1, 2, \dots, \nu)$$

von der Determinante 1 in die Form ψ übergehen. Wenn wir auf f zuerst die Substitution (r) und darauf die Substitution (τ) anwenden, so erhalten wir zuerst die Form φ und dann die Form ψ . Zu derselben Form ψ müssen wir nun offenbar gelangen, indem wir auf f unmittelbar die zusammengesetzte Substitution $(r) \cdot (\tau)$, d. i. die Substitution

$$(s): x_i = \sum_{k=1}^{\nu} \left(\sum_{h=1}^{\nu} r_i^h \tau_h^k \right) \eta_k \quad (i = 1, 2, \dots, n)$$

anwenden. Hiernach ist die Form ψ durch f vermittle der Substitution

$$(s): x_i = \sum_{k=1}^{\nu} s_i^k \eta_k \quad (i = 1, 2, \dots, n)$$

darstellbar, in welcher

$$(42) \quad s_i^k = \sum_{h=1}^{\nu} r_i^h \tau_h^k \quad \left(\begin{array}{l} k = 1, 2, \dots, \nu \\ i = 1, 2, \dots, n \end{array} \right)$$

ist.

Nach einem bekannten Satz gelten die Beziehungen

$$|s_i^k| = |r_i^h| \cdot |\tau_h^k|, \quad \left(\begin{array}{l} i = i_1, i_2, \dots, i_{\nu} \\ k, h = 1, 2, \dots, \nu \end{array} \right)$$

d. i.

$$s \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_{\nu} \end{pmatrix} = r \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_{\nu} \end{pmatrix}.$$

Diese Beziehungen zeigen uns, daß der größte Teiler der sämtlichen aus ν Reihen der Substitution (s) gebildeten Unterdeterminanten gleich dem größten Teiler der sämtlichen aus ν Reihen der Substitution (r) gebildeten Unterdeterminanten ist. Folglich ist die Darstellung (s) eine eigentliche, sobald (r) eine eigentliche Darstellung ist.

Wir wollen die Darstellung $(s) = (r) \cdot (\tau)$ der Form ψ der Darstellung (r) der Form φ äquivalent nennen und schreiben $(r_{\varphi}) \sim (s_{\psi})$.

Unter den ν -reihigen Determinanten der Substitution (r) gibt es mindestens eine von Null verschiedene; es sei etwa

$$r \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_{\nu} \end{pmatrix} \neq 0.$$

Wählen wir unter den Gleichungen (42) diejenigen ν aus, welche den Werten $i = i_1, i_2, \dots, i_{\nu}$ und $k = k$ entsprechen und lösen sie nach den Koeffizienten τ_h^k auf, so erhalten wir

$$(43) \quad \tau_h^k = \sum_i \frac{1}{r \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_{\nu} \end{pmatrix}} \cdot \frac{\partial r \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_{\nu} \end{pmatrix}}{\partial r_i^h} \cdot s_i^k \quad (i = i_1, i_2, \dots, i_{\nu})$$

Also können die Koeffizienten τ_h^k vermittle der Größen r_i^h und s_i^k ausgedrückt werden. Daraus schließen wir, daß zwei verschiedene Transformationen (τ) von φ in ψ niemals dieselbe einer gegebenen Darstellung (r_{φ}) äquivalente Darstellung (s_{ψ}) liefern können.

II. Wenn zwei Formen φ und ψ durch f vermittle zweier Substitutionen (r) und (s) dargestellt sind, welche die Bedingungen

$$(44) \quad r \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_{\nu} \end{pmatrix} = s \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_{\nu} \end{pmatrix} \quad (i = 1, 2, \dots, n)$$

erfüllen, so ist stets $\varphi \sim \psi$ und $(r_{\varphi}) \sim (s_{\psi})$.

In der Tat: sind die Zahlen $r_i^k (k > \nu)$ so gewählt, daß die Substitution

$$R: x_i = \sum_{k=1}^n r_i^k \xi_k \quad (i = 1, 2, \dots, n)$$

die Determinante 1 besitzt, und setzen wir $s_i^k = r_i^k (k > \nu)$, so zeigen die Gleichungen (44) unmittelbar, daß die Substitution

$$S: x_i = \sum_{k=1}^n s_i^k \eta_k \quad (i = 1, 2, \dots, n)$$

gleichfalls von der Determinante 1 ist. Die beiden Formen $\bar{R} \cdot f \cdot R = \Phi$ und $\bar{S} \cdot f \cdot S = \Psi$ können wir schreiben

$$\Phi = \sum_{i,k=1}^n \alpha_{ik} \xi_i \xi_k \quad \text{und} \quad \Psi = \sum_{i,k=1}^n \beta_{ik} \eta_i \eta_k;$$

dann wird

$$\varphi = \sum_{i,k=1}^{\nu} \alpha_{ik} \xi_i \xi_k \quad \text{und} \quad \psi = \sum_{i,k=1}^{\nu} \beta_{ik} \eta_i \eta_k.$$

Offenbar gilt

$$\Psi = \bar{R}^{-1} S \cdot \Phi \cdot R^{-1} S.$$

Man erkennt also, daß sich vermittels der Substitution

$$R^{-1} S: \xi_i = \sum_{k=1}^n \tau_i^k \eta_k, \quad (i = 1, 2, \dots, n)$$

in der

$$\tau_i^k = \sum_{h=1}^n \frac{\partial |R|}{\partial r_h^i} \cdot s_h^k \quad (i, k = 1, 2, \dots, n)$$

ist, die Form Φ in Ψ verwandelt.

Wie man leicht einsieht, lassen sich die Zahlen $\frac{\partial |R|}{\partial r_h^i} (i > \nu)$ mit Hilfe der Größen $r \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_\nu \end{pmatrix}$ und mittels Zahlen $r_{h_0}^k (k > \nu)$ ausdrücken. Folglich gilt

$$\frac{\partial |R|}{\partial r_h^i} = \frac{\partial |S|}{\partial s_h^i} \quad (i > \nu),$$

also

$$\tau_i^k = \begin{cases} 1 & (i = k) \\ 0 & (i \neq k) \end{cases} \quad (i > \nu),$$

und die Substitution $R^{-1} S$ nimmt die Form an

$$R^{-1} S: \xi_i = \sum_{k=1}^n \tau_i^k \eta_k \quad (i \leq \nu), \quad \xi_i = \eta_i \quad (i > \nu).$$

Wir finden sonach

$$\beta_{lm} = \sum_{i,k=1}^v \alpha_{ik} \tau_i^l \tau_k^m, \quad (l, m = 1, 2, \dots, v)$$

und die Form φ geht vermittle der Substitution

$$(\tau): \quad \xi_i = \sum_{k=1}^v \tau_i^k \eta_k \quad (i = 1, 2, \dots, v)$$

in ψ über. Die Determinante dieser Substitution ist gleich der Determinante der Substitution $R^{-1}S$, d. i. gleich der Einheit. Mithin ist die Form φ der Form ψ äquivalent.

Setzen wir die Systeme R und $R^{-1}S$ zusammen, so müssen wir zu dem System S gelangen. Daraus ergeben sich die Gleichungen

$$s_i^k = \sum_{h=1}^v r_i^h \tau_h^k, \quad (k = 1, 2, \dots, v; i = 1, 2, \dots, n)$$

aus denen ersichtlich ist, daß die Darstellung (s_ψ) der Darstellung (r_φ) äquivalent ist.

III. Mit Hilfe des Satzes II. können wir leicht den folgenden Satz beweisen:

P. Ist $\varphi \sim \psi'$, $\varphi \sim \psi''$ und $(r_\varphi) \sim (s_{\psi'})$, $(r_\varphi) \sim (s_{\psi''})$, so ist $(s_{\psi'}) \sim (s_{\psi''})$.

Wir fassen die sämtlichen Darstellungen einer Form φ , welche einer gegebenen Darstellung dieser Form äquivalent sind, in eine *Gruppe von Darstellungen* der Form φ zusammen. Aus dem Satze P. folgt alsdann, daß zwei Darstellungen (r_φ) derselben Gruppe stets äquivalent, zwei Darstellungen (r_φ) aus verschiedenen Gruppen nicht äquivalent sind.

Aus den Gleichungen (42) und (43) erkennen wir, daß die Anzahl aller verschiedenen Darstellungen (r) einer Form φ , welche in einer Gruppe von Darstellungen dieser Form auftreten [d. i. die Dichtigkeit einer derartigen Gruppe (r_φ)] gleich der Anzahl der sämtlichen verschiedenen Transformationen (von der Determinante 1) der Form φ in sich selbst ist. Daraus schließen wir leicht, daß die Darstellungsgruppen zweier äquivalenter Formen gleiche Dichtigkeit besitzen.

Wir nennen zwei *Gruppen* von Darstellungen (r_φ) und (s_ψ) *äquivalent*, wenn irgendeine Darstellung der einen Gruppe irgendeiner Darstellung der andern äquivalent ist. Der Satz P. zeigt dann, daß zwei beliebige Darstellungen nur dann äquivalent sind, wenn sie äquivalenten Gruppen von Darstellungen angehören. Hiernach bilden die sämtlichen Darstellungen (s_ψ) , welche einer gegebenen Darstellung (r_φ) äquivalent sind, eine Gruppe von Darstellungen der Form ψ .

Kap. XV. Adjungierte Darstellungen und adjungierte Darstellungsgruppen.

I. Es möge eine primitive Form f von n Variablen mit Hilfe einer Substitution

$$R: x_i = \sum_{k=1}^n r_i^k \xi_k \quad (i = 1, 2, \dots, n)$$

von der Determinante 1 in eine äquivalente Form $\Phi = \{\alpha_{ik}\}$ ($i, k = 1, 2, \dots, n$)

übergehen, und es möge φ die Form $\sum_{i,k=1}^v \alpha_{ik} \xi_i \xi_k$ von v Variablen bezeichnen. Es sei f' die adjungierte Form von f ,

$$R': x'_i = \sum_{k=1}^n r_i'^k \xi'_k \quad (i = 1, 2, \dots, n)$$

die adjungierte Substitution von R und $\Phi' = \{\alpha'_{ik}\}$ ($i, k = 1, 2, \dots, n$) die adjungierte Form von Φ . In Kap. XII haben wir gesehen, daß die Form f' mit Hilfe der Substitution R' in Φ' übergeht. Wir setzen die

Form $\sum_{i,k=1}^{n-v} \alpha'_{ik} \xi'_i \xi'_k$ von $n - v = v'$ Variablen gleich φ' . Ist die Determinante der Form φ gleich $\sigma_v d_{v-1}(\varphi)$ und die Determinante von φ' gleich $\sigma_{v'} d'_{v'-1}(\varphi')$, so gilt die Beziehung

$$(\varphi) = (\varphi') \cdot (-1)^{I(\varphi)}.$$

Wir wollen sagen, die Darstellung

$$(r'): x'_i = \sum_{k=1}^{v'} r_i'^k \xi'_k \quad (i = 1, 2, \dots, n)$$

der Form φ' durch die Form f' sei der Darstellung

$$(r): x_i = \sum_{k=1}^v r_i^k \xi_k \quad (i = 1, 2, \dots, n)$$

der Form φ durch die Form f adjungiert, und wollen uns des Zeichens $(r_\varphi) \times (r'_{\varphi'})$ bedienen. — Aus der Reziprozität zwischen den Formen f und f' erhellt, daß, wenn die Darstellung $(r'_{\varphi'})$ der Darstellung (r_φ) adjungiert ist, umgekehrt die Darstellung (r_φ) der Darstellung $(r'_{\varphi'})$ adjungiert ist.

Da die Systeme R und R' adjungiert sind, bestehen für die Darstellungen (r) und (r') die sämtlichen Gleichungen

$$(45) \quad r(1, 2, \dots, v) = r'(1, 2, \dots, v') \\ (i_1, i_2, \dots, i_v) = (i'_1, i'_2, \dots, i'_{v'}),$$

in denen die Indizes i und i' so gewählt sein sollen, daß die n Zahlen i_h ; $n+1-i'_h$, abgesehen von der Reihenfolge den n Zahlen $1, 2, \dots, n$ gleich sind und die Permutation

$$(i_1, i_2, \dots, i_\nu, n+1-i'_\nu, \dots, n+1-i'_2, n+1-i'_1)$$

aus $(1, 2, \dots, n)$ vermittelt einer geraden Anzahl von Transpositionen hervorgeht.

II. Umgekehrt sind die beiden Darstellungen (r_φ) und $(r'_{\varphi'})$ stets adjungiert, sobald sie die sämtlichen Bedingungen (45) erfüllen.

In der Tat: seien die Zahlen $\varrho_i'^k$ ($k > \nu'$) so gewählt, daß die Substitution

$$(\varrho'): x_i' = \sum_{k=1}^{\nu'} r_i'^k \xi_k' + \sum_{k=\nu'+1}^n \varrho_i'^k \xi_k' \quad (i = 1, 2, \dots, n)$$

die Determinante 1 besitzt, und sei

$$(\varrho): x_i = \sum_{k=1}^{\nu} \varrho_i^k \xi_k + \sum_{k=\nu+1}^n R_i^k \xi_k \quad (i = 1, 2, \dots, n)$$

die adjungierte Substitution von (ϱ') . Es wird dann

$$\varrho \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_\nu \end{pmatrix} = r' \begin{pmatrix} 1, 2, \dots, \nu' \\ i_1', i_2', \dots, i_{\nu'}' \end{pmatrix},$$

woraus sich wegen der Gleichungen (45)

$$r \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_\nu \end{pmatrix} = \varrho \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_\nu \end{pmatrix}$$

ergibt. Wir erkennen nunmehr, daß die Substitution

$$R: x_i = \sum_{k=1}^{\nu} r_i^k \xi_k + \sum_{k=\nu+1}^n R_i^k \xi_k \quad (i = 1, 2, \dots, n)$$

die Determinante 1 besitzt und daß die adjungierte Substitution von R die Form erhält:

$$R': x_i' = \sum_{k=1}^{\nu'} r_i'^k \xi_k' + \sum_{k=\nu'+1}^n R_i'^k \xi_k' \quad (i = 1, 2, \dots, n).$$

Also ist die Darstellung (r') in der Tat der Darstellung (r) adjungiert.

III. Wir können jetzt den folgenden Satz beweisen:

Ist $(r_\varphi) \sim (s_\psi)$ ($\varphi \sim \psi$) und $(r_\varphi) \times (r'_{\varphi'})$, $(s_\psi) \times (s'_{\psi'})$, so ist stets $(r'_{\varphi'}) \sim (s'_{\psi'})$ ($\varphi' \sim \psi'$) und $(r_\varphi) \times (s'_{\psi'})$, $(s_\psi) \times (r'_{\varphi'})$.

In der Tat ergibt die Voraussetzung

$$r \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_\nu \end{pmatrix} = s \begin{pmatrix} 1, 2, \dots, \nu \\ i_1, i_2, \dots, i_\nu \end{pmatrix} = r' \begin{pmatrix} 1, 2, \dots, \nu' \\ i_1', i_2', \dots, i_{\nu'}' \end{pmatrix} = s' \begin{pmatrix} 1, 2, \dots, \nu' \\ i_1', i_2', \dots, i_{\nu'}' \end{pmatrix},$$

woraus unmittelbar die Richtigkeit der aufgestellten Behauptung folgt.

Dieser Satz zeigt, daß zwei äquivalenten Darstellungen stets dieselben adjungierten Darstellungen zukommen.

Wir nennen zwei *Gruppen* von Darstellungen (r) , (r') *adjungiert*, sobald irgendeine Darstellung der einen Gruppe irgendeiner Darstellung der andern adjungiert ist. Aus dem vorstehenden Satz schließen wir, daß, wenn zwei Gruppen von Darstellungen adjungiert sind, jede Darstellung der einen Gruppe jeder Darstellung der andern adjungiert ist.

Den sämtlichen äquivalenten Darstellungsgruppen, welche zu Formen einer bestimmten Klasse φ von ν Variablen gehören, sind äquivalente Darstellungsgruppen einer bestimmten Klasse φ' von ν' Variablen adjungiert. Die beiden Klassen φ und φ' besitzen hiernach eine gewisse Reziprozität in bezug auf die Formen f und f' , und man kann sehr bemerkenswerte Relationen zwischen den Indizes, den Ordnungen und den Genera dieser beiden Klassen aufstellen. Wir leiten an dieser Stelle nur die Relation zwischen den Indizes her.

Q. Bezeichnen wir durch J den Index von φ , durch J' den Index von φ' und durch I den Index von f und f' , so hat die Gleichung statt

$$J + J' = I. \quad (J \leq I, J' \leq I)$$

Denn setzen wir die aus den ersten $h(=1, 2, \dots, n)$ Horizontal- und Vertikalreihen von Φ resp. Φ' gebildeten Unterdeterminanten gleich $\sigma_h d_{h-1} \Delta_h$ resp. $\sigma'_h d'_{h-1} \Delta'_h$, so ist die Zahl J resp. J' gleich der Anzahl der sämtlichen negativen Größen aus der Reihe $\frac{\Delta_h}{\Delta_{h-1}} (h=1, 2, \dots, \nu; \Delta_0=1)$ resp. aus der Reihe $\frac{\Delta'_h}{\Delta'_{h-1}} (h=1, 2, \dots, \nu'; \Delta'_0=1)$, während die Zahl I gleich der Anzahl der negativen Größen aus der Reihe $\frac{\Delta_h}{\Delta_{h-1}} (h=1, 2, \dots, n)$ oder aus der Reihe $\frac{\Delta'_h}{\Delta'_{h-1}} (h=1, 2, \dots, n)$ wird. Nun haben wir nach Kap. XII die Relationen $\Delta_h = (-1)^I \cdot \Delta'_{n-h} (h=0, 1, 2, \dots, n)$; dieselben führen mit Leichtigkeit zu dem angegebenen Resultate.

Von besonderem Interesse ist der Fall, in welchem die Form φ eine Ordnung

$$\begin{pmatrix} \sigma_1, \sigma_2, \dots, \sigma_{\nu-2}, \sigma_{\nu-1} \\ o_1, o_2, \dots, o_{\nu-2}, \sigma_{\nu} o_{\nu-1} \cdot m \end{pmatrix}, \quad J$$

und die Form φ' eine Ordnung

$$\begin{pmatrix} \sigma'_1, \sigma'_2, \dots, \sigma'_{\nu'-2}, \sigma'_{\nu'-1} \\ o'_1, o'_2, \dots, o'_{\nu'-2}, \sigma'_{\nu'} o'_{\nu'-1} \cdot m \end{pmatrix}, \quad J'$$

besitzt.

Kap. XVI. Darstellung von ganzen Zahlen durch Formen mit n Variablen.

Wir wollen weiterhin insbesondere den Fall betrachten, in welchem eine der Zahlen ν, ν' gleich 1 ist. Es sei $\nu = 1, \nu' = n - 1$.

Wenn die Form f mit Hilfe der Substitution

$$(t): x_i = t_i \xi \quad (i = 1, 2, \dots, n)$$

in eine (einvариablige) Form $b\xi^2$ übergeht, so wird die Zahl b durch f vermittels der Zahlen

$$(t): x_i = t_i$$

dargestellt, und umgekehrt ist die Form $b\xi^2$ stets durch f darstellbar, sobald die Zahl b durch f dargestellt werden kann.

Wir nennen die Darstellung $x_i = t_i$ der Zahl b durch die Form f eigentlich, wenn die Darstellung von $b\xi^2$ durch f eigentlich ist, d. h. wenn der größte gemeinsame Teiler τ der n Zahlen t_i gleich 1 ist.

So oft der Teiler τ größer als 1 ist, muß b den Faktor $\tau^2 > 1$ enthalten, und die Darstellung $x_i = t_i$ der Zahl b durch f liefert die eigentliche Darstellung $x_i = \frac{t_i}{\tau}$ der Zahl $\frac{b}{\tau^2}$ durch f . Wir gelangen infolgedessen zu allen überhaupt möglichen Darstellungen einer Zahl b durch eine Form f , indem wir die sämtlichen quadratischen Divisoren τ^2 von b aufsuchen und alle eigentlichen Darstellungen der Zahlen $\frac{b}{\tau^2}$ durch die Form f bestimmen.

Wir setzen die Form f als primitiv voraus. Zwei eigentliche Darstellungen (t) und (t^0) einer Form $b\xi^2$ oder einer Zahl b durch die Form f werden nach unseren Definitionen nur dann äquivalent sein, wenn die sämtlichen Gleichungen $t_i = t_i^0$ ($i = 1, 2, \dots, n$) statthaben, d. h. wenn sie identisch sind. Infolgedessen sprechen sich die in Kap. XV aufgestellten Sätze für den Fall $\nu = 1$ folgendermaßen aus*):

I. Zu jeder eigentlichen Darstellung (t) einer Zahl b durch die Form f sind eigentliche Darstellungen (r') gewisser Formen φ' von $n - 1$ Variablen und der Determinante $(-1)^I d'_{n-2} \cdot b$ durch die Form f' adjungiert.

I'. Zu jeder eigentlichen Darstellung (r') einer Form φ' von $n - 1$ Variablen und der Determinante $(-1)^I d'_{n-2} \cdot b$ durch die Form f' ist eine einzige eigentliche Darstellung (t) der Zahl b durch die Form f adjungiert.

II. Zu zwei äquivalenten Darstellungen (r') und (s') zweier äquivalenter Formen φ' und ψ' von $n - 1$ Variablen und der Determinante $(-1)^I d'_{n-2} \cdot b$ durch die Form f' ist eine und dieselbe Darstellung (t) der Zahl b durch die Form f adjungiert.

*) Siehe Gauß, Disquisitiones arithmeticae, art. 280.

II'. Wenn zu einer und derselben Darstellung (t) der Zahl b durch die Form f zwei Darstellungen (r') und (s') zweier Formen φ' und ψ' mit $n - 1$ Variablen und von der Determinante $(-1)^I d'_{n-2} \cdot b$ durch die Form f' adjungiert sind, so sind die Formen φ' und ψ' und die Darstellungen (r') und (s') äquivalent.

Um die sämtlichen eigentlichen Darstellungen einer Zahl b durch eine primitive Form f zu bestimmen, können wir jetzt in folgender Weise verfahren: wir wählen aus einer jeden Formenklasse mit $n - 1$ Variablen und von der Determinante $(-1)^I d'_{n-2} \cdot b$ einen Repräsentanten φ' aus, suchen die sämtlichen nicht-äquivalenten Gruppen von eigentlichen Darstellungen dieser Formen φ' durch die Form f' auf und bestimmen zu jeder dieser nicht-äquivalenten Gruppen die einzige adjungierte eigentliche Darstellung der Zahl b durch die Form f . Auf diese Weise wird man zu den sämtlichen überhaupt möglichen eigentlichen Darstellungen der Zahl b durch die Form f gelangen und zwar zu einer jeden dieser Darstellungen nur ein einziges Mal.

Kap. XVII. Darstellungen von Formen mit $n - 1$ Variablen durch Formen mit n Variablen.

Wir wenden uns jetzt zu einer näheren Untersuchung der eigentlichen Darstellungen einer Form mit $n - 1$ Variablen durch eine primitive Form mit n Variablen.

I. Es möge eine Form $\varphi' = \sum_{i,k=1}^{n-1} b'_{ik} \xi'_i \xi'_k$ von der Determinante $(-1)^I d'_{n-2} \cdot b$ durch eine primitive Form $f' = \sum_{i,k=1}^n a'_{ik} x'_i x'_k$ mittels einer Substitution

$$(\vartheta'): x'_i = \sum_{k=1}^{n-1} \vartheta'_i{}^k \xi'_k \quad (i = 1, 2, \dots, n)$$

[[eigentlich]] dargestellt werden.

Bestimmen wir n Zahlen t'_i , so daß die Substitution

$$(t'): x'_i = \sum_{k=1}^{n-1} \vartheta'_i{}^k \xi'_k + t'_i \xi'_n \quad (i = 1, 2, \dots, n)$$

die Determinante 1 besitzt, so wird die Form $\overline{(t)} \cdot f' \cdot (t) = B'$ der Form f' äquivalent sein und sich schreiben lassen

$$B' = \sum_{i,k=1}^{n-1} b'_{ik} \xi'_i \xi'_k + 2 \sum_{i=1}^{n-1} b'_i \xi'_i \xi'_n + b'_n \xi'^2_n.$$

Wir bezeichnen durch $f = \sum_{i,k=1}^n a_{ik} x_i x_k$ die zu f' adjungierte Form, durch

$$(t): x_i = t_i \xi + \sum_{k=1}^{n-1} \vartheta_i^k \xi_k \quad (i = 1, 2, \dots, n)$$

die adjungierte Substitution von (t') und durch

$$B = b \xi^2 + 2 \sum_{i=1}^{n-1} b_i \xi \xi_i + \sum_{i,k=1}^{n-1} b_{ik} \xi_i \xi_k$$

die zu B' adjungierte Form. Die n Zahlen t_i drücken sich dabei durch die gegebenen Koeffizienten ϑ_i^k aus. Es gelten die Beziehungen

$$(46) \quad \sum_{i=1}^n t_i' t_{n-i+1} = 1, \quad \sum_{h=1}^{n-1} \vartheta_{n-i+1}^{n-h} \vartheta_k^h + t_{n-i+1}' t_k = \begin{cases} 1 & (i=k) \\ 0 & (i \neq k) \end{cases},$$

und wir bekommen $B = \overline{(t)} \cdot f \cdot (t)$, d. i.

$$b_{im} = \sum_{i,k=1}^n a_{ik} \vartheta_i^i \vartheta_k^m$$

und

$$(47) \quad b = \sum_{i,k=1}^n a_{ik} t_i t_k, \quad b_h = \sum_{i,k=1}^n a_{ik} t_i \vartheta_k^h.$$

Mit Benutzung der Summen

$$T_i = \sum_{h=1}^n a_{ih} t_h = a_{i1} t_1 + a_{i2} t_2 + \dots + a_{in} t_n = \frac{1}{2} \frac{\partial b}{\partial t_i}$$

können wir die Formeln (47) auch schreiben

$$(48) \quad b = \sum_{i=1}^n T_i t_i, \quad b_h = \sum_{i=1}^n T_i \vartheta_i^h.$$

Wir setzen $|b_{ik}'| = |\varphi'|$, $|b_{ik}| = |\varphi|$ und

$$\frac{\partial |\varphi'|}{\partial b_{ik}'} = (-1)^I \cdot d'_{n-3} \cdot c_{n-i, n-k}, \quad \frac{\partial |\varphi|}{\partial b_{ik}} = (-1)^I \cdot d_{n-3} \cdot c'_{n-i, n-k}.$$

Da die Formen f und f' adjungiert sind, gelten nach einem bekannten Determinantensatze die Gleichungen

$$\begin{aligned} \{(-1)^I \cdot d'_{n-2} b\} \{(-1)^I \cdot d'_{n-2} b_{ik}\} &= \{(-1)^I \cdot d'_{n-2} b_i\} \{(-1)^I \cdot d'_{n-2} b_k\} \\ &= (-1)^I \cdot d'_{n-1} \cdot \{(-1)^I \cdot d'_{n-3} \cdot c_{ik}\} \end{aligned}$$

oder

$$(49) \quad -o_1 c_{ik} = b_i b_k - b b_{ik},$$

aus denen wir die Kongruenzen gewinnen

$$(50) \quad -o_1 c_{ik} \equiv b_i b_k \pmod{b}$$

oder

$$-o_1 \cdot \left(\sum_{i,k=1}^{n-1} c_{ik} \xi_i \xi_k \right) \equiv \left(\sum_{i=1}^{n-1} b_i \xi_i \right)^2 \pmod{b}.$$

Jede Darstellung (ϑ') der Form φ' durch die Form f' liefert auf solche Weise bestimmte Lösungen $(b_1, b_2, \dots, b_{n-1})$ der $\frac{n(n-1)}{2}$ Kongruenzen (50). Wir sagen, die Darstellung (ϑ') gehöre zu diesen Lösungen (b_h) ($h = 1, 2, \dots, n-1$)*).

II. Wenn wir den n Zahlen $t'_1, t'_2, \dots, t'_n$ alle möglichen Werte erteilen, welche die Bedingung $\sum_{i=1}^n t'_i t_{n-i+1} = 1$ erfüllen, so erhalten wir

alle Lösungen (b_h) der Kongruenzen (50), zu welchen die gegebene Darstellung (ϑ') der Form φ' gehört. Es besteht nun der Satz:

Alle Lösungen $(b_1, b_2, \dots, b_{n-1})$ der Kongruenzen (50), zu welchen die nämliche Darstellung (ϑ') der Form φ' gehört, sind nach dem Modul b kongruent.

In der Tat, die Summen

$$\sum_{h=1}^{n-1} \vartheta'^{n-h}_{n-i+1} \cdot b_h = \vartheta'^{n-1}_{n-i+1} \cdot b_1 + \vartheta'^{n-2}_{n-i+1} \cdot b_2 + \dots + \vartheta'^1_{n-i+1} \cdot b_{n-1}$$

nehmen mit Hilfe der Formeln (48) und (46) die Werte an

$$\sum_{h=1}^{n-1} \left(\vartheta'^{n-h}_{n-i+1} \cdot \sum_{k=1}^n \vartheta^h_k T_k \right) = \sum_{k=1}^n \left(T_k \sum_{h=1}^{n-1} \vartheta^h_k \vartheta'^{n-h}_{n-i+1} \right) = -t'_{n-i+1} \sum_{k=1}^n T_k t_k + T_i.$$

Es kommt also

$$(51) \quad \sum_{h=1}^{n-1} \vartheta'^{n-h}_{n-i+1} \cdot b_h - \sum_{k=1}^n a_{ik} t_k = -b t'_{n-i+1} \equiv 0 \pmod{b}.$$

$$(i = 1, 2, \dots, n)$$

Fassen wir irgendwelche $n-1$ von diesen n Gleichungen zusammen, etwa alle diejenigen, welche einem Index $i \neq g$ entsprechen, so können wir dieselben nach den $n-1$ Größen b_h auflösen, und es ergibt sich

$$t_g b_h - \sum_{(i)} \left(\frac{\partial t_g}{\partial \vartheta'^{n-h}_{n-i+1}} \cdot \sum_{k=1}^n a_{ik} t_k \right) = b \cdot \vartheta^h_g \equiv 0 \pmod{b} \quad (i \neq g).$$

* Siehe Gauß, Disquisitiones arithmeticae, art. 282.

Es ist klar, daß die Größen $t_k, \frac{\partial t_g}{\partial \vartheta'^{n-i+1}}$ mittels der Koeffizienten $\vartheta_i'^k$ der Darstellung (ϑ') ausgedrückt werden können; es sind daher auch die Reste der Zahlen

$$t_1 b_h, t_2 b_h, \dots, t_n b_h$$

nach dem Modul b durch diese Zahlen $\vartheta_i'^k$ vollständig bestimmt. Da die Darstellung (ϑ') eine eigentliche ist, können die n Zahlen $t_1, t_2, \dots, t_n$ keinen gemeinsamen Teiler größer als 1 besitzen, und es müssen sich daher unter diesen n Zahlen solche befinden, die zu einem beliebigen Primfaktor q von b relativ prim sind. Infolge dieses Umstandes sind auch die Reste der $n-1$ Zahlen b_h für jede in b aufgehende Primzahlpotenz q^i und mithin für den Modul b selbst eindeutig durch die Koeffizienten $\vartheta_i'^k$ bestimmt, und hieraus geht unmittelbar das Behauptete hervor.

Setzen wir jetzt voraus, man könne n Zahlen t_i' $\left(\sum_{i=1}^n t_i' t_{n-i+1} = 1\right)$

so finden, daß die Darstellung (ϑ') zu der Wurzel (b_h) der Kongruenzen (50) gehört, und es möge $(\tilde{b}_h)$ eine andere, nach dem Modul b der Wurzel (b_h) kongruente Wurzel dieser Kongruenzen sein. Alsdann kann man n Zahlen

$\tilde{t}_i' \left(\sum_{i=1}^n \tilde{t}_i' t_{n-i+1} = 1\right)$ derart finden, daß die Darstellung (ϑ') zu dieser Wurzel $(\tilde{b}_h)$ gehört.

In der Tat, es möge $\tilde{b}_h = b_h + b e_h$ sein. Für die Zahlen $\tilde{t}_i'$ müssen die Beziehungen

$$(51) \quad \sum_{h=1}^{n-1} \vartheta'^{n-h}_{n-i+1} \cdot \tilde{b}_h - \sum_{k=1}^n a_{ik} t_k = -b \cdot \tilde{t}'_{n-i+1}$$

statthaben. Die Differenz der Gleichungen (51) und $(\tilde{51})$ ergibt sogleich

$$(52) \quad \tilde{t}_i' = t_i' - \sum_{h=1}^{n-1} \vartheta_i'^{n-h} e_h.$$

Daraus geht hervor, daß die Bestimmung der Zahlen $\tilde{t}_i'$ jedenfalls nur auf eine einzige Weise möglich sein kann. Führen wir nun für die Zahlen $\tilde{t}_i'$ die Ausdrücke (52) ein, so ist die Gleichung $(\tilde{51})$ wirklich erfüllt. Multiplizieren wir dann diese Gleichung mit t_i und bilden die Summe über alle Werte $i = 1, 2, \dots, n$, so bekommen wir

$$b \cdot \sum_{i=1}^n \tilde{t}'_{n-i+1} t_i = \sum_{i,k=1}^n a_{ik} t_i t_k,$$

d. i.

$$\sum_{i=1}^n \tilde{t}'_{n-i+1} t_i = 1.$$

Die Darstellung (ϑ') gehört also in der Tat zu der Wurzel $(\tilde{b}_h)$ der Kongruenzen (50).

Ist insbesondere $e_h = 0$, $\tilde{b}_h = b_h$ ($h = 1, 2, \dots, n-1$), so wird $\tilde{t}'_i = t'_i$ ($i = 1, 2, \dots, n$). Man sieht demnach, daß es nicht zwei verschiedene Substitutionen (t') geben kann, welche dieselben Koeffizienten $\vartheta_i'^k$ besitzen und die Form f' durch dieselbe Form B' ersetzen.

III. Zu den eben bewiesenen Sätzen gelangen wir auch auf folgendem Wege*):

Es seien $\tilde{t}'_1, \tilde{t}'_2, \dots, \tilde{t}'_n$ irgendwelche Zahlen, welche ebenso wie die Zahlen $t'_1, t'_2, \dots, t'_n$ der Gleichung

$$\sum_{i=1}^n \tilde{t}'_i \cdot t_{n-i+1} = 1$$

genügen, und es möge die Form f' durch die Substitution

$$(\tilde{t}'): \quad x'_i = \sum_{k=1}^{n-1} \vartheta_i'^k \tilde{\xi}'_k + \tilde{t}'_i \tilde{\xi}' \quad (i = 1, 2, \dots, n)$$

in eine Form

$$\tilde{B}' = \sum_{i,k=1}^{n-1} b'_{ik} \tilde{\xi}'_i \tilde{\xi}'_k + 2 \sum_{i=1}^{n-1} \tilde{b}'_i \tilde{\xi}'_i \tilde{\xi}' + \tilde{b}' \tilde{\xi}'^2$$

übergehen. Der Substitution $(\tilde{t}')$ sei die Substitution

$$(\tilde{t}): \quad x_i = t_i \tilde{\xi} + \sum_{k=1}^{n-1} \tilde{\vartheta}_i^k \tilde{\xi}_k \quad (i = 1, 2, \dots, n)$$

und der Form $\tilde{B}'$ die Form

$$\tilde{B} = b \tilde{\xi}^2 + 2 \sum_{i=1}^{n-1} \tilde{b}_i \tilde{\xi} \tilde{\xi}_i + \sum_{i,k=1}^{n-1} \tilde{b}_{ik} \tilde{\xi}_i \tilde{\xi}_k$$

adjungiert. Wie man aus Kap. XIV (II) ersieht, geht alsdann die Form B mit Hilfe der Substitution

$$(\tilde{t})^{-1} \cdot (\tilde{t}): \quad \xi = \tilde{\xi} + \sum_{h=1}^{n-1} e_h \cdot \tilde{\xi}_h, \quad \xi_h = \tilde{\xi}_h \quad (h = 1, 2, \dots, n-1),$$

in welcher

$$e_h = \sum_{i=1}^n t'_{n-i+1} \vartheta_i^h$$

*) Siehe Gauß, Disquisitiones arithmeticae, art. 282.

ist, in die Form $\tilde{B}$ über. Wir bekommen daher

$$\tilde{b}_h = b_h + b e_h, \quad \tilde{b}_h \equiv b_h \pmod{b},$$

und die beiden Lösungen (b_h) und $(\tilde{b}_h)$ der Kongruenzen (50) sind in der Tat kongruent modulo b .

Die Substitution $(t')^{-1} \cdot (\tilde{t}')$, welche die Form B' in $\tilde{B}'$ verwandelt, läßt sich jetzt schreiben

$$(t')^{-1} \cdot (\tilde{t}'): \quad \xi'_h = \tilde{\xi}'_h - e_{n-h} \cdot \tilde{\xi}' \quad (h = 1, 2, \dots, n-1), \quad \xi' = \tilde{\xi}'.$$

Durch Zusammensetzung der beiden Systeme (t') und $(t')^{-1}(\tilde{t}')$ müssen wir zu dem System $(\tilde{t}')$ gelangen. Auf diese Weise erhalten wir von neuem die Bedingungen (52). Führen wir aber für die Zahlen $\tilde{t}'_i$ die Werte (52) ein, so ist die Substitution $(\tilde{t}')$ in der Tat von der Determinante 1 und verwandelt die Form f' in eine Form $\tilde{B}'$, deren adjungierte Form $\tilde{B}$ anstelle der Koeffizienten b_h die Zahlen $\tilde{b}_h$ aufweist.

IV. Man kann leicht die folgende Relation beweisen, welche später Anwendung finden wird:

$$(53) \quad \prod_{h=1}^{n-1} o_h^2 = -(n-2)bb' \cdot \prod_{h=1}^{n-1} o_h + o_1 o_1' \cdot \sum_{i,k=1}^{n-1} c_{ik} c'_{n-i, n-k}.$$

In der Tat hat man

$$|\varphi| = \sum_{k=1}^{n-1} b_{ik} \frac{\partial |\varphi|}{\partial b_{ik}},$$

d. i.

$$(54) \quad o_1 o_2 \dots o_{n-2} b' = \sum_{k=1}^{n-1} b_{ik} \cdot c'_{n-i, n-k}.$$

Die Determinante der Form B läßt sich schreiben

$$(-1)^I \cdot d_{n-1} = b \cdot |\varphi| - \sum_{i,k=1}^{n-1} b_i b_k \frac{\partial |\varphi|}{\partial b_{ik}};$$

daraus ergibt sich

$$\prod_{h=1}^{n-1} o_h^2 = bb' \cdot \prod_{h=1}^{n-1} o_h - o_1' \cdot \sum_{i,k=1}^{n-1} b_i b_k \cdot c'_{n-i, n-k}.$$

Führen wir hierin statt der Zahlen $b_i b_k$ die Zahlen $-o_1 c_{ik} + b b_{ik}$ ein und benutzen die Beziehung (54), so bekommen wir sofort die behauptete Gleichung.

V. Es ist klar, daß es überhaupt keine eigentlichen Darstellungen der Form $\varphi' = \{b'_{ik}\}$ von der Determinante $(-1)^I \cdot d'_{n-2} \cdot b$ durch die Form f' geben kann, sobald nicht die Größen $c_{n-i, n-k} = (-1)^I \cdot \frac{1}{d'_{n-3}} \cdot \frac{\partial |\varphi'|}{\partial b'_{ik}}$ ganze Zahlen werden. Sind aber diese sämtlichen Größen ganze Zahlen,

so wird eine jede eigentliche Darstellung (ϑ') der Form φ' durch f' zu einer einzigen Lösung $(b_h) \pmod{b}$ der Kongruenzen

$$-o_1 c_{ik} \equiv b_i b_k \pmod{b}$$

gehören, und wir werden sonach alle möglichen eigentlichen Darstellungen (ϑ') der Form φ' durch f' bekommen, indem wir die sämtlichen modulo b inkongruenten Systeme

$$(b_1, b_2, \dots, b_{n-1})$$

aufsuchen, welche diese Kongruenzen erfüllen, und für jedes dieser Systeme die sämtlichen eigentlichen Darstellungen (ϑ') bestimmen, welche zu demselben gehören.

Für ein bestimmtes gegebenes System $(b_h) \pmod{b}$ kann dies auf folgende Weise geschehen:

Wir bezeichnen die ganzen Zahlen

$$\frac{o_1 c_{ik} + b_i b_k}{b}$$

durch b_{ik} und untersuchen die Form

$$B = b \xi^2 + 2 \sum_{i=1}^{n-1} b_i \xi \xi_i + \sum_{i,k=1}^{n-1} b_{ik} \xi_i \xi_k.$$

Diese Form geht mit Hilfe der Substitution

$$\xi = \eta - \sum_{h=1}^{n-1} \frac{b_h}{b} \eta_h, \quad \xi_h = \eta_h \quad (h = 1, 2, \dots, n-1)$$

in eine Form

$$B_0 = b \eta^2 + \sum_{i,k=1}^{n-1} \frac{o_1 c_{ik}}{b} \eta_i \eta_k$$

über. Wir ersehen hieraus, daß die Determinante von B gleich $(-1)^I \cdot d_{n-1}$ ist. Ferner gilt $\frac{\partial |B|}{\partial b_{n-i, n-k}} = (-1)^I \cdot d_{n-2} \cdot b'_{ik}$, und wir setzen

$$\frac{\partial |B|}{\partial b} = (-1)^I \cdot d_{n-2} \cdot b' \quad \text{und} \quad \frac{\partial |B|}{\partial b_{n-i}} = (-1)^I \cdot d_{n-2} \cdot b'_i.$$

Aus den vorhergehenden Sätzen leuchtet ein, daß wenn die Form B der Form f nicht äquivalent ist, keine eigentlichen, zu der Wurzel (b_h) der Kongruenzen (50) gehörenden Darstellungen der Form φ' durch f' existieren. Sobald aber $B \sim f$ ist, so wird auch die zu B adjungierte Form B' der Form f' äquivalent sein und sich schreiben lassen

$$B' = \sum_{i,k=1}^{n-1} b'_{ik} \xi'_i \xi'_k + 2 \sum_{i=1}^{n-1} b'_i \xi'_i \xi' + b' \xi'^2.$$

Eine jede Substitution

$$(t'): \quad x'_i = \sum_{k=1}^{n-1} \vartheta'_i{}^k \xi'_k + t'_i \xi' \quad (i=1, 2, \dots, n)$$

von der Determinante 1, mit deren Hilfe die Form f' in B' übergeht, liefert dann eine eigentliche Darstellung

$$(\vartheta'): \quad x'_i = \sum_{k=1}^{n-1} \vartheta'_i{}^k \xi'_k \quad (i=1, 2, \dots, n)$$

von φ' durch f' , und wir gelangen, indem wir sämtliche verschiedenen Substitutionen (t') von der Determinante 1 bilden, durch welche f' in B' transformiert wird, zu allen überhaupt möglichen Darstellungen (ϑ'_i) , welche zu der Wurzel (b_h) gehören, und zwar zu einer jeden dieser Darstellungen ein einziges Mal. Denn wir haben gesehen, daß zwei verschiedene Transformationen (t') niemals die gleichen Koeffizienten $\vartheta'_i{}^k$ darbieten können.

Kap. XVIII. Index, Ordnung und Genus der durch eine Form von n Variablen darstellbaren Formen von $n-1$ Variablen.

Es sei eine primitive Form $f = \sum_{i,k=1}^n a_{ik} x_i x_k$ von einem Index I , einer

Ordnung $O: \begin{pmatrix} \sigma_h \\ o_h \end{pmatrix}$ und einem Genus G gegeben, und es möge die Zahl b mittels eines Systems

$$(t): \quad x_i = t_i$$

durch f dargestellt sein.

Da die Koeffizienten a_{ii} , $2a_{ik}$ sämtlich durch σ_1 teilbar sind, wird die Zahl b den Faktor σ_1 enthalten. Es bedeute δ das Vorzeichen der Zahl b und m den absoluten Wert von $\frac{b}{\sigma_1}$. Dann wird $b = \delta \cdot \sigma_1 m$. Wir betrachten insbesondere Zahlen b , für welche die Größe m zu der Determinante Δ von f relativ prim ist.

Wir nennen den größten gemeinsamen Teiler T der n Zahlen $T_i = \sum_{k=1}^n a_{ik} t_k$ den *Teiler der Darstellung* (t) , und sagen, eine Darstellung (t) sei *primär*, wenn ihr Teiler T gleich 1 ist. Eine primäre Darstellung ist stets eigentlich; denn der größte gemeinsame Teiler τ der n Zahlen t_k geht in allen Summen T_i und folglich auch in der Zahl T auf. Ist daher $T=1$, so wird auch der Teiler τ der Einheit gleich sein. — Aus den Gleichungen (48) erhellt, daß der Teiler T der Darstellung (t) in den sämtlichen Zahlen $b, b_1, b_2, \dots, b_{n-1}$ aufgehen wird. Andererseits erkennen

wir aus den Gleichungen (51), daß der größte Teiler der Zahlen b, b_h in den sämtlichen Zahlen T_i aufgeht. Es stimmt demnach der Teiler T der Darstellung (t) mit dem größten Teiler der Zahlen b, b_h überein. Die Kongruenzen $b \equiv 0, b_h \equiv 0 \pmod{T}$ ergeben unmittelbar $\Delta \equiv 0 \pmod{T}$. Folglich wird, falls die Zahl m zu Δ relativ prim sein soll, der Teiler T in σ_1 aufgehen, und eine [[eigentliche]] Darstellung (t) der Zahl b wird stets primär sein, außer in dem Falle, daß $\sigma_1 = 2, m \equiv 1 \pmod{2}$ ist und die Zahlen T_i alle gerade ausfallen.

Der Darstellung (t) der Zahl b durch die Form f ist eine Darstellung (ϑ') einer Form $\varphi' = \sum_{i,k=1}^{n-1} b'_{ik} \xi'_i \xi'_k$ von der Determinante $(-1)^I d'_{n-2} \cdot b$ durch die Form $f' = \sum_{i,k=1}^n a'_{ik} x'_i x'_k$ adjungiert. Wir wollen die Beziehungen untersuchen, welche zwischen den Indizes, den Ordnungen und den Genera der Form φ' und der Form f' bestehen.

Index der Form φ' .

Nach dem Satze Q. (Kap. XV) muß der Index der Form φ' gleich I oder $I-1$ sein, je nachdem die Zahl b positiv ($\delta = 1$) oder negativ ($\delta = -1$) ist. Da der Index einer Form von $n-1$ Variablen stets zwischen den Grenzen 0 und $n-1$ eingeschlossen ist, wird die Form φ' niemals durch die Form f' darstellbar sein, wenn $I = 0$ und $b < 0$ oder $I = n$ und $b > 0$ ist.

Ordnung der Form φ' .

Es sei der größte gemeinsame Teiler der sämtlichen Koeffizienten b'_{ik} der Form φ' gleich e' , und es seien $e'_1, e'_2, \dots, e'_{n-2}$ die Invarianten σ und $\tau'_1, \tau'_2, \dots, \tau'_{n-2}$ die Invarianten σ der primitiven Form $\frac{\varphi'}{e'}$. Es gilt alsdann die Beziehung

$$(55) \quad \sigma_1 d'_{n-2} m = e'^{n-1} \cdot e'_1{}^{n-2} e'_2{}^{n-3} \dots e'_{n-2}.$$

Wir bezeichnen durch $q', q^{h'}$ die höchsten Potenzen einer Primzahl q , welche in den Größen e', e'_h aufgehen. Wir wollen jetzt die Zahlen $\varepsilon', \varepsilon'_h$ durch die Zahlen ω'_h ausdrücken.

1. Es bedeute zunächst q eine in m aufgehende Primzahl. Wäre die Größe e' oder eine der $n-3$ ersten Invarianten $e'_1, e'_2, \dots, e'_{n-3}$ durch diese Primzahl q teilbar, so müßten die sämtlichen ganzen Zahlen $c_{ik} = (-1)^I \cdot \frac{1}{d'_{n-3}} \cdot \frac{\partial |\varphi'|}{\partial b'_{n-i, n-k}}$ gleichfalls durch q teilbar sein, und die Gleichung (53) gäbe $\prod_{h=1}^{n-1} \omega'_h$ oder $\Delta \equiv 0 \pmod{q}$, was gegen unsere Voraus-

setzung streitet, daß die Zahl m zu Δ relativ prim ist. Wir finden sonach für jede in m aufgehende Primzahl $\varepsilon' = 0$; $\varepsilon_1' = 0$, $\varepsilon_2' = 0$, ..., $\varepsilon_{n-3}' = 0$. Wegen (55) wird dann die Größe $q^{t'_{n-2}}$ der größten Potenz von q gleich sein, welche in $\sigma_1 m$ oder in b aufgeht. Indem wir dieses Resultat für die sämtlichen in m enthaltenen Primzahlen q anwenden, erkennen wir, daß die Invariante e'_{n-2} durch m teilbar sein muß.

2. Es bedeute jetzt q eine ungerade Primzahl p , welche nicht in m aufgeht. Da die Zahl b alsdann zu p relativ prim ist, so können wir die Zahlen $(b_1, b_2, \dots, b_{n-1})$, zu welchen die Darstellung (ϑ') gehört, so wählen, daß sie neben den Kongruenzen (50) für den Modul b noch den weiteren Kongruenzen

$$b_1 \equiv 0, b_2 \equiv 0, \dots, b_{n-1} \equiv 0 \pmod{p^f}$$

für irgend einen Modul $p^t (> p^{v_n-1(p)})$ genügen. Es wird alsdann

$$o_1 c_{ik} \equiv b b_{ik} \pmod{p^t},$$

und die Form B besitzt für den Modul p^t einen Rest

$$B \equiv \begin{pmatrix} b, & 0, & \dots, & 0 \\ 0, & \frac{o_1 c_{11}}{b}, & \dots, & \frac{o_1 c_{1, n-1}}{b} \\ \cdot & \cdot & \cdot & \cdot \\ 0, & \frac{o_1 c_{n-1, 1}}{b}, & \dots, & \frac{o_1 c_{n-1, n-1}}{b} \end{pmatrix} \pmod{p^t}.$$

Schlüsse, welche denen von Kap. III ganz analog sind, zeigen jetzt, daß die Form $\sum_{i, k=1}^{n-1} c_{ik} \xi_i \xi_k$ primitiv in bezug auf p sein muß und daß die höchsten in den Invarianten o dieser Form aufgehenden Potenzen von p bzw. gleich $p^{\omega'_{n-2}}, p^{\omega'_{n-3}}, \dots, p^{\omega'_1}$ sind. Andererseits müssen diese Potenzen gleich $p^{\varepsilon'_{n-2}}, p^{\varepsilon'_{n-3}}, \dots, p^{\varepsilon'_1}$ sein; denn die Form $\{c_{ik}\}$ ist ein Multiplum der zu $\frac{\vartheta'}{e'}$ adjungierten Form. Wir gewinnen also die Beziehungen

$$\varepsilon_1' = \omega_1', \varepsilon_2' = \omega_2', \dots, \varepsilon_{n-2}' = \omega_{n-2}'.$$

Indem wir jetzt diese Werte der Zahlen ε'_k in die Formel (55) einführen, finden wir noch $\varepsilon' = 0$. Die Form ϑ' muß also in bezug auf p primitiv sein.

3. Es möge endlich die Primzahl q gleich 2 sein. Gemäß unserer Annahme, daß die Zahl m zu Δ relativ prim sei, werden wir die beiden Fälle $m \equiv 1 \pmod{2}$ und $m \equiv 0, \Delta \equiv 1 \pmod{2}$ zu untersuchen haben.

I. Wir betrachten zunächst den Fall $m \equiv 1 \pmod{2}$.

($\sigma_1 = 1$). Ist in diesem Falle $\sigma_1 = 1$, so wird die Zahl b ungerade sein, und wir können infolgedessen die Zahlen $(b_1, b_2, \dots, b_{n-1})$, zu welchen die Darstellung (ϑ') gehört, so wählen, daß sie neben den Kongruenzen (50) nach dem Modul b noch den Kongruenzen

$$b_1 \equiv 0, b_2 \equiv 0, \dots, b_{n-1} \equiv 0 \pmod{2^t}$$

für irgend einen Modul $2^t (> \sigma_{n-1} \cdot 2^{e_{n-1}(2)})$ genügen. Alsdann wird

$$o_1 c_{ik} \equiv b b_{ik} \pmod{2^t}$$

und folglich

$$B \equiv \begin{pmatrix} b, & 0 & , \dots, & 0 \\ 0, & \frac{o_1 c_{11}}{b} & , \dots, & \frac{o_1 c_{1, n-1}}{b} \\ \cdot & \cdot & \cdot & \cdot & \cdot \\ 0, & \frac{o_1 c_{n-1, 1}}{b} & , \dots, & \frac{o_1 c_{n-1, n-1}}{b} \end{pmatrix} \pmod{2^t}.$$

Aus diesem Rest von B ersehen wir, daß die Form $\sum_{i, k=1}^{n-1} c_{ik} \xi_i \xi_k$ zu 2 primitiv ausfällt, daß die Zahlen $\varepsilon', \varepsilon'_h$ durch die Gleichungen

$$(56) \quad \varepsilon' = 0, \varepsilon'_1 = \omega'_1, \varepsilon'_2 = \omega'_2, \dots, \varepsilon'_{n-2} = \omega'_{n-2}$$

gegeben sind und daß die Potenzen $\sigma'_{n-h} \cdot 2^{\partial_{h-1}}$ mit den kleineren der Potenzen

$$\tau'_{n-h} \cdot 2^{\partial_{h-1}}, \quad \tau'_{n-h-1} \cdot 2^{\partial_h}$$

übereinstimmen oder, was auf dasselbe hinauskommt, daß die Invarianten σ'_{n-h} mit den kleineren der je zwei Zahlen

$$\tau'_{n-h} \quad \text{und} \quad \tau'_{n-h-1} \cdot 2^{\omega_1 + \omega_2 + \dots + \omega_h}$$

übereinstimmen.

Wir wollen nun annehmen, daß die $\kappa_1 - 1$ ($1 \leq \kappa_1 \leq n$) ersten der Größen ω_h , nämlich $\omega_1, \omega_2, \dots, \omega_{\kappa_1-1}$ gleich Null seien, während die κ_1^{te} dieser Zahlen, ω_{κ_1} , von Null verschieden sei. Alsdann ist nach Kap. IV $\sigma_1 = 1, \sigma_2 = 1, \dots, \sigma_{\kappa_1-1} = 1; \sigma_{\kappa_1} = 1$. Die Größe $\tau'_{n-h-1} \cdot 2^{\omega_1 + \omega_2 + \dots + \omega_h}$ wird für $h < \kappa_1$ gleich τ'_{n-h-1} und für $h \geq \kappa_1$ größer oder gleich $2\tau'_{n-h-1} \geq \tau'_{n-h}$. Für $h \geq \kappa_1$ ist daher immer $\sigma'_{n-h} = \tau'_{n-h}$.

Die Relationen (56) ergeben

$$\varepsilon'_{n-\kappa_1} \geq 1, \varepsilon'_{n-\kappa_1+1} = 0, \varepsilon'_{n-\kappa_1+2} = 0, \dots, \varepsilon'_{n-2} = 0.$$

Die $\kappa_1 - 2$ letzten Invarianten τ'_h können infolgedessen nach den in Kap. IV gegebenen Sätzen nur die Werte

$$(57_1) \quad \tau'_{n-\kappa_1+1} = 1, \tau'_{n-\kappa_1+2} = 1, \dots, \tau'_{n-2} = 1$$

oder die Werte

$$(57_2) \quad \tau'_{n-\kappa_1+1} = 2, \tau'_{n-\kappa_1+2} = 1, \dots, \tau'_{n-2} = 2$$

annehmen.

Der letztere Fall (57₂) ist an die Bedingungen $\kappa_1 - 1 \equiv 0 \pmod{2}$ und $\kappa_1 - 1 > 0$ gebunden. Wenn also $\kappa_1 \equiv 0 \pmod{2}$ und auch wenn $\kappa_1 = 1$ wird, erhalten wir stets $\tau'_h = \sigma'_h$ ($h = 1, 2, \dots, n-2$). Ist dagegen $\kappa_1 \equiv 1 \pmod{2}$ und $\kappa_1 > 1$, so sind die Fälle (57₁) und (57₂) alle beide möglich.

Der zweite dieser Fälle tritt offenbar nur dann ein, wenn die Invariante τ'_{n-2} gleich 2 ist, d. h. wenn die Zahlen

$$c_{ii} = \frac{1}{o_1} (b b_{ii} - b_i^2) \equiv b_{ii} - b_i \pmod{2}$$

alle kongruent 0 (mod 2) sind, oder, was auf das nämliche hinauskommt, wenn die $n - 1$ Kongruenzen

$$(58) \quad b_i \equiv b_{ii} \pmod{2} \quad [b \equiv 1 \pmod{2}]$$

gelten.

Wir können voraussetzen, daß der Rest $f \pmod{2}$ von der in Kap. III angegebenen Gestalt $f_{(x_1)}$ ist. Geht dann f durch die Substitution

$$(t): \quad x_i = t_i \xi + \sum_{k=1}^{n-1} \vartheta_i^k \xi_k \quad (i = 1, 2, \dots, n)$$

in B über, so wird

$$b_i \equiv t_1 \vartheta_1^i + t_2 \vartheta_2^i + \dots + t_{x_1} \vartheta_{x_1}^i, \quad b_{ii} \equiv \vartheta_1^i + \vartheta_2^i + \dots + \vartheta_{x_1}^i \pmod{2},$$

und der Fall (57₂) ist durch die Bedingungen

$$(59) \quad (t_1 - 1) \vartheta_1^i + (t_2 - 1) \vartheta_2^i + \dots + (t_{x_1} - 1) \vartheta_{x_1}^i \equiv 0 \pmod{2} \quad (i = 1, 2, \dots, n-1)$$

gekennzeichnet. Zu diesen Kongruenzen können wir noch die Kongruenz

$$(60) \quad (t_1 - 1)t_1 + (t_2 - 1)t_2 + \dots + (t_{x_1} - 1)t_{x_1} \equiv 0 \pmod{2},$$

die wegen $t_h(t_h - 1) \equiv 0 \pmod{2}$ evident ist, hinzufügen. Da die Determinante der Substitution (t) ungerade ($\equiv 1$) ist, können die x_1 -reihigen Unterdeterminanten des Systems

$$|t_h, \vartheta_h^1, \vartheta_h^2, \dots, \vartheta_h^{n-1}| \quad (h = 1, 2, \dots, x_1)$$

nicht sämtlich gerade sein. Demnach gibt es unter den n Kongruenzen (59), (60) x_1 , deren Determinante ungerade ist. Lösen wir dann diese x_1 Kongruenzen nach den x_1 Größen $t_1 - 1, t_2 - 1, \dots, t_{x_1} - 1$ auf, so bekommen wir

$$t_1 - 1 \equiv 0, t_2 - 1 \equiv 0, \dots, t_{x_1} - 1 \equiv 0 \pmod{2}.$$

Man erkennt also, daß die Kongruenzen (58) die einzige Lösung

$$t_1 \equiv 1, t_2 \equiv 1, \dots, t_{x_1} \equiv 1 \pmod{2}$$

zulassen.

($\sigma_1 = 2$.) Es sei jetzt $\sigma_1 = 2$ und mithin $b = \delta \cdot \sigma_1 m \equiv 2 \pmod{4}$. Die $n - 1$ Zahlen $(b_1, b_2, \dots, b_{n-1})$ sind dann entweder alle gerade, oder es befinden sich unter ihnen auch ungerade Größen. Im ersteren Falle wird $T_i = \sum_{k=1}^n a_{ik} t_k \equiv 0 \pmod{2} \quad (i = 1, 2, \dots, n)$, und folglich ist die gegebene Darstellung (t) der Zahl b nur im zweiten Falle primär.

Wir wollen annehmen, von den Größen ω_h seien die $x_1 - 1$ ($1 \leq x_1 \leq n$) ersten, nämlich $\omega_1, \omega_2, \dots, \omega_{x_1-1}$, gleich Null, während die x_1^{te} dieser

Zahlen, ω_{x_1} , von Null verschieden sein möge, und wir wollen annehmen, daß der Rest $f \pmod{4}$ von der in Kap. III angegebenen Gestalt $f_{(x_1)}$ sei. Dann gelten die Kongruenzen

$$T_1 \equiv t_2, T_2 \equiv t_1, T_3 \equiv t_4, T_4 \equiv t_3, \dots, T_{x_1-1} \equiv t_{x_1}, T_{x_1} \equiv t_{x_1-1};$$

$$T_h \equiv 0 \pmod{2} \quad (h > x_1),$$

und es werden die Zahlen T_i nur dann sämtlich gerade ausfallen, wenn

$$t_1 \equiv 0, t_2 \equiv 0, \dots, t_{x_1} \equiv 0 \pmod{2}$$

wird. Diese Kongruenzen bilden demnach die notwendige und hinreichende Bedingung dafür, daß die $n-1$ Zahlen b_h sämtlich durch 2 teilbar werden. Diese Kongruenzen sind jedoch mit der Annahme $b \equiv 2 \pmod{4}$ nur in dem Falle verträglich, daß $2^{\omega_{x_1}} \sigma_{x_1+1} = 2$ und also $2^{\omega_{x_1}} = 2, \sigma_{x_1+1} = 1$ wird. In der Tat, ist $2^{\omega_{x_1}} \sigma_{x_1+1} \equiv 0 \pmod{4}$, so sieht man leicht, daß in dem Rest $f \equiv \Phi_{(1)} + 2^{\omega_{x_1}} f_{(x_1)} \pmod{4}$ alle Koeffizienten $2^{\omega_{x_1}} a_{i_i}^{(x_1)}, 2^{\omega_{x_1}} a_{i_k}^{(x_1)}$ kongruent Null

(mod 4) werden. Infolgedessen ist eine jede Zahl b , welche durch f mittels gerader Zahlen $t_1, t_2, \dots, t_{x_1}$ dargestellt wird, durch 4 teilbar.

Wenn von den $n-1$ Zahlen $b_1, b_2, \dots, b_{n-1}$ einige ungerade ausfallen, so werden sich auch unter den Zahlen $c_{ii} = \frac{1}{o_1} (-b_i^2 + b b_{ii})$ ungerade

Größen befinden. Die Form $\Phi = \sum_{i,k=1}^{n-1} c_{ik} \xi_i \xi_k$ wird daher in bezug auf 2 primitiv ausfallen und sich in einen Repräsentanten

$$\tilde{\Phi} = \begin{pmatrix} c^*, & c_1^*, & \dots, & c_{n-2}^* \\ c_1^*, & c_{11}^*, & \dots, & c_{1,n-2}^* \\ \dots & \dots & \dots & \dots \\ c_{n-2}^*, & c_{n-2,1}^*, & \dots, & c_{n-2,n-2}^* \end{pmatrix}$$

verwandeln lassen, in welchem $c^* \equiv 1 \pmod{2}$ ist und die Koeffizienten $c_1^*, \dots, c_{n-2}^*$ kongruent Null nach einem Modul $2^t (> \sigma_{n-1} \cdot 2^{e_{n-1}(2)})$ sind.

Bedeutet M den größten Teiler der sämtlichen Koeffizienten von Φ , so wird die Form $\frac{\delta \cdot \Phi}{M}$ der Form $\frac{\varphi'}{e}$ adjungiert sein. Wir wollen durch $\frac{\tilde{\varphi}'}{e}$ die

zu $\frac{\delta \cdot \tilde{\Phi}}{M}$ adjungierte Form bezeichnen. Dann ist die Form $\tilde{\varphi}'$ der Form φ' äquivalent. Anstatt der vorliegenden Darstellung (ϑ') der Form φ' durch f' können wir jetzt irgendeine äquivalente Darstellung $(\tilde{\vartheta}')$ der Form $\tilde{\varphi}'$ durch f' betrachten. Für diese Darstellung $(\tilde{\vartheta}')$ möge anstelle der Form B eine Form

$$\tilde{B} = b \xi^2 + 2 \sum_{i=0}^{n-2} \tilde{b}_i \tilde{\xi}_i \tilde{\xi}_i + \sum_{i,k=0}^{n-2} \tilde{b}_{ik} \tilde{\xi}_i \tilde{\xi}_k$$

treten.

Wir erhalten die Gleichungen

$$-o_1 c^* = \tilde{b}_0^2 - b \cdot \tilde{b}_{00}, \quad -o_1 c_i^* = \tilde{b}_0 \tilde{b}_i - b \cdot \tilde{b}_{0i}, \quad -o_1 c_{ik}^* = \tilde{b}_i \tilde{b}_k - b \cdot \tilde{b}_{ik}.$$

Da wir $b \equiv 0 \pmod{2}$ und $c^* \equiv 1$, $c_i^* \equiv 0 \pmod{2}$ vorausgesetzt haben, ergeben sich die Kongruenzen

$$\tilde{b}_0 \equiv 1; \quad \tilde{b}_1 \equiv 0, \dots, \tilde{b}_{n-2} \equiv 0 \pmod{2}.$$

Die Darstellung $(\tilde{\vartheta}')$ bestimmt nur die Reste der Zahlen $\tilde{b}_i$ nach dem Modul $b \equiv 2 \pmod{4}$. Wir können daher die Zahlen so wählen, daß sie den Kongruenzen

$$\tilde{b}_1 \equiv c_1^*, \dots, \tilde{b}_{n-2} \equiv c_{n-2}^* \pmod{2^{t+1}}$$

genügen. Alsdann werden wegen $\tilde{b}_0 \tilde{b}_i - b \cdot \tilde{b}_{0i} = -o_1 c_i^*$ die Kongruenzen $\tilde{b}_{0i} \equiv 0 \pmod{2^t}$ und wegen $\tilde{b}_i \tilde{b}_k - b \cdot \tilde{b}_{ik} = -o_1 c_{ik}^*$ die Kongruenzen $c_{ii}^* \equiv 0$, $2c_{ik}^* \equiv 0 \pmod{4}$ statthaben, und die Form $\tilde{B}$ wird für den Modul 2^t einen Rest

$$\tilde{B} \equiv \begin{pmatrix} b, & b_0 \\ b_0, & b_{00} \\ & \frac{o_1 c_{11}^*}{b}, & \dots, & \frac{o_1 c_{1, n-2}^*}{b} \\ & \dots & \dots & \dots \\ & \frac{o_1 c_{n-2, 1}^*}{b}, & \dots, & \frac{o_1 c_{n-2, n-2}^*}{b} \end{pmatrix} \pmod{2^t}$$

liefern, in welchem $b \equiv 0$, $b_0 \equiv 1$, $b_{00} \equiv 0$, $\frac{o_1 c_{ii}^*}{b} \equiv 0 \pmod{2}$ ist.

Dieser Rest von $\tilde{B}$ zeigt, daß die höchste in allen Koeffizienten c_{ik}^* der Form $\sum_{i,k=1}^{n-2} c_{ik}^* \xi_i \xi_k$ aufgehende Potenz von 2 gleich $2^{\omega'_{n-2}+1}$ ist, sowie daß für die in bezug auf 2 primitive Form $\left\{ \frac{c_{ik}^*}{2^{\omega'_{n-2}+1}} \right\}$ die Invarianten $2^{\omega^{(2)}}$ gleich $2^{\omega'_{n-3}}, \dots, 2^{\omega'_1}$ und die Invarianten σ gleich $\sigma'_{n-3}, \dots, \sigma'_1$ sind. Indem wir dieses Resultat auf den Rest der Form $\tilde{\Phi} \pmod{2^t}$ anwenden, erkennen wir, daß die Invarianten $2^{\omega^{(2)}}$ dieser Form gleich $\sigma_1 \cdot 2^{\omega'_{n-2}}, 2^{\omega'_{n-3}}, \dots, 2^{\omega'_1}$ und die Invarianten σ dieser Form gleich $\sigma'_{n-2}, \sigma'_{n-3}, \dots, \sigma'_1$ sind. Andererseits ist offenbar, daß die Invarianten $2^{\omega^{(2)}}$ der Form $\tilde{\Phi}$ mit den Zahlen $2^{\varepsilon'_{n-2}}, 2^{\varepsilon'_{n-3}}, \dots, 2^{\varepsilon'_1}$ und die Invarianten σ dieser Form mit den Größen $\tau'_{n-2}, \tau'_{n-3}, \dots, \tau'_1$ übereinstimmen. Folglich wird

$$\begin{aligned} \varepsilon'_1 &= \omega'_1, \quad \varepsilon'_2 = \omega'_2, \quad \dots, \quad \varepsilon'_{n-3} = \omega'_{n-3}, \quad 2^{\varepsilon'_{n-2}} = \sigma_1 \cdot 2^{\omega'_{n-2}}, \\ \tau'_1 &= \sigma'_1, \quad \tau'_2 = \sigma'_2, \quad \dots, \quad \tau'_{n-3} = \sigma'_{n-3}, \quad \tau'_{n-2} = \sigma'_{n-2} = 1. \end{aligned}$$

Führen wir diese Werte der Zahlen ε'_h in die Relation (55) ein, so finden wir noch $\varepsilon' = 0$.

II. Wir betrachten endlich den Fall, in welchem $m \equiv 0 \pmod{2}$ und $\Delta \equiv 1 \pmod{2}$ ist. Es sei 2^n die höchste in m aufgehende Potenz von 2. Wir sahen in 1., daß die Zahlen $\varepsilon', \varepsilon'_h$ durch die Gleichungen

$$\varepsilon' = 0, \varepsilon'_1 = 0, \varepsilon'_2 = 0, \dots, \varepsilon'_{n-3} = 0, 2^{\varepsilon'_{n-2}} = \sigma_1 \cdot 2^n$$

gegeben sind.

Wir haben sonach nur noch die Invarianten τ'_h zu bestimmen. Die letzte dieser Invarianten, τ'_{n-2} , wird, da $2^{\varepsilon'_{n-2}} \geq 2$ ist, stets gleich 1. Die übrigen $n-3$ Invarianten τ'_h müssen nach den in Kap. IV gegebenen Sätzen entweder die Werte

$$\tau'_1 = 1, \tau'_2 = 1, \dots, \tau'_{n-3} = 1$$

oder die Werte

$$\tau'_1 = 2, \tau'_2 = 1, \dots, \tau'_{n-3} = 2$$

erhalten. Der letztere dieser beiden Fälle ist an die Bedingung $n-2 \equiv 0 \pmod{2}$, also $n \equiv 0 \pmod{2}$ gebunden. Außerdem muß, wie wir leicht einsehen, eine jede Invariante τ'_h die entsprechende Invariante σ'_h als Faktor enthalten. Wir gewinnen daher das folgende Resultat: Wenn $n \equiv 1 \pmod{2}$ wird, in welchem Falle stets $\sigma'_1 = 1, \sigma'_2 = 1, \dots, \sigma'_{n-2} = 1, \sigma'_{n-1} = 1$ ist, so haben wir $\tau'_1 = 1, \tau'_2 = 1, \dots, \tau'_{n-3} = 1, \tau'_{n-2} = 1$; wenn dagegen $n \equiv 0 \pmod{2}$ ist, so wird, falls $\sigma'_1 = 1, \sigma'_2 = 1, \dots, \sigma'_{n-2} = 1, \sigma'_{n-1} = 1$ ist, entweder $\tau'_1 = 1, \tau'_2 = 1, \dots, \tau'_{n-3} = 1, \tau'_{n-2} = 1$ oder $\tau'_1 = 2, \tau'_2 = 1, \dots, \tau'_{n-3} = 2, \tau'_{n-2} = 1$, und falls $\sigma'_1 = 2, \sigma'_2 = 1, \dots, \sigma'_{n-2} = 1, \sigma'_{n-1} = 2$ wird, haben wir stets $\tau'_1 = 2, \tau'_2 = 1, \dots, \tau'_{n-3} = 2, \tau'_{n-2} = 1$.

Wir können die Sätze, zu denen wir gelangt sind, in der folgenden Weise zusammenfassen:

Wenn eine primäre Darstellung (t) einer Zahl $b = \delta \cdot \sigma_1 m$ (m relativ prim zu Δ) durch die Form f einer Darstellung (ϑ') einer Form φ' von $n-1$ Variablen und der Determinante $(-1)^I \cdot d'_{n-2} b$ durch die Form f' adjungiert ist, so fällt φ' primitiv aus und gehört zu einer Ordnung

$$\begin{pmatrix} \tau'_h \\ e'_h \end{pmatrix} (h = 1, 2, \dots, n-2), \quad J' = I - \frac{1-\delta}{2},$$

deren Invarianten e'_h den Gleichungen

$$e'_1 = o'_1, e'_2 = o'_2, \dots, e'_{n-3} = o'_{n-3}, e'_{n-2} = o'_{n-2} \cdot \sigma_1 m$$

genügen und deren Invarianten τ'_h entweder die Gleichungen

$$(I) \quad \tau'_1 = \sigma'_1, \tau'_2 = \sigma'_2, \dots, \tau'_{n-3} = \sigma'_{n-3}, \tau'_{n-2} = \sigma'_{n-2}$$

oder auch, falls $\sigma_1 = 1, m \equiv \kappa_1 \pmod{2}, \kappa_1 > 1$ ist, die Gleichungen

$$(II) \quad \begin{cases} \tau'_1 = \sigma'_1, & \tau'_2 = \sigma'_2, \dots, \tau'_{n-\kappa_1} = \sigma'_{n-\kappa_1}, \\ \tau'_{n-\kappa_1+1} = 2, \tau'_{n-\kappa_1+2} = 1, \dots, \tau'_{n-3} = 1, \tau'_{n-2} = 2, \\ \sigma'_{n-\kappa_1+1} = 1, \sigma'_{n-\kappa_1+2} = 1, \dots, \sigma'_{n-3} = 1, \sigma'_{n-2} = 1, \end{cases}$$

oder, falls $\kappa_1 = n$, $\sigma_1 = 1$, $m \equiv \kappa_1 \equiv 0 \pmod{2}$, $\kappa_1 > 2$ ist, die Gleichungen

$$(II) \quad \begin{cases} \tau'_1 = 2, \tau'_2 = 1, \dots, \tau'_{n-3} = 2, \tau'_{n-2} = 1, \\ \sigma'_1 = 1, \sigma'_2 = 1, \dots, \sigma'_{n-3} = 1, \sigma'_{n-2} = 1 \end{cases}$$

erfüllen.

Wir wollen die Ordnung der Form φ' , je nachdem die Gleichungen (I) oder (II) statthaben, durch $O_I'(b)$ oder durch $O_{II}'(b)$ bezeichnen.

Sei jetzt φ die der Form φ' adjungierte Form. Aus dem obigen

Satz erhellt, daß φ gleich $\sum_{i,k=1}^{n-1} c_{ik} \xi_i \xi_k$ oder gleich $-\sum_{i,k=1}^{n-1} c_{ik} \xi_i \xi_k$ sein wird, je nachdem die Zahl b positiv oder negativ ist. Wir können daher

$\varphi = \delta \cdot \sum_{i,k=1}^{n-1} c_{ik} \xi_i \xi_k$ setzen, und die Form φ wird primitiv sein und der Ordnung

$$\left(\begin{matrix} \tau'_{n-2}, \tau'_{n-3}, \dots, \tau'_1 \\ e'_{n-2}, e'_{n-3}, \dots, e'_1 \end{matrix} \right), \quad J', \quad [e'_{n-2} = o'_{n-2} \cdot \sigma_1 m]$$

angehören.

Wir wählen die Form φ' in ihrer Klasse so aus, daß φ einen Hauptrest für den Modul b vorstellt. Dann wird

$$\delta \cdot \varphi = \begin{pmatrix} c^*, & c_1^*, & c_2^*, & \dots, & c_{n-2}^* \\ c_1^*, & c_{11}^*, & c_{12}^*, & \dots, & c_{1,n-2}^* \\ c_2^*, & c_{21}^*, & c_{22}^*, & \dots, & c_{2,n-2}^* \\ \dots & \dots & \dots & \dots & \dots \\ c_{n-2}^*, & c_{n-2,1}^*, & c_{n-2,2}^*, & \dots, & c_{n-2,n-2}^* \end{pmatrix} \equiv \begin{pmatrix} c^*, & 0, & 0, & \dots, & 0 \\ 0, & 0, & 0, & \dots, & 0 \\ 0, & 0, & 0, & \dots, & 0 \\ \dots & \dots & \dots & \dots & \dots \\ 0, & 0, & 0, & \dots, & 0 \end{pmatrix} \pmod{b},$$

wo c^* zu b relativ prim ist, und den $\frac{n(n-1)}{2}$ Kongruenzen

$$(61) \quad -o_1 c^* \equiv \tilde{b}_0^2 \pmod{b}, \quad -o_1 c_i^* \equiv \tilde{b}_0 \tilde{b}_i \pmod{b}, \quad -o_1 c_{ik}^* \equiv \tilde{b}_i \tilde{b}_k \pmod{b} \\ (i, k = 1, 2, \dots, n-2)$$

kann nicht anders genügt werden, als wenn man

$$(62) \quad -o_1 c^* \equiv \tilde{b}_0^2 \pmod{b}$$

und

$$\tilde{b}_1 \equiv 0, \tilde{b}_2 \equiv 0, \dots, \tilde{b}_{n-2} \equiv 0 \pmod{b}$$

hat. Es wird daher die Anzahl der sämtlichen inkongruenten Systeme

$(\tilde{b}_0, \tilde{b}_1, \dots, \tilde{b}_{n-2}) \pmod{b}$, welche den $\frac{n(n-1)}{2}$ Kongruenzen (61) genügen,

mit der Anzahl der inkongruenten Lösungen $\tilde{b}_0 \pmod{b}$ der einzigen Kongruenz (62) übereinstimmen. Diese letztere Anzahl ist, sobald die Kongruenz (62) überhaupt keine Lösungen besitzt, gleich Null, dagegen wenn die Kongruenz (62) lösbar ist und wenn in der Zahl b im ganzen μ ungerade Primzahlen p aufgehen, gleich 2^μ , falls $b \equiv 1 \pmod{2}$ oder $b \equiv 2 \pmod{4}$ ist, oder gleich $2^{\mu+1}$, wenn $b \equiv 4 \pmod{8}$ ist, oder gleich $2^{\mu+2}$, falls $b \equiv 0 \pmod{8}$ ist.

Genus der Form φ' .

Wir wollen jetzt zeigen, daß die Form φ' nur einem einzigen Genus $G_I'(b)$ oder $G_{II}'(b)$ der Ordnung $O_I'(b)$ resp. $O_{II}'(b)$ angehören kann und daß die Charaktere dieses Genus völlig durch die Charaktere des Genus G' der Form f' bestimmt sind.

Vergleichen wir zu diesem Zweck die Charaktere der Form φ' und der Form B' in bezug auf einen Modul q^t miteinander. Der Einfachheit halber denken wir uns die Form φ' als eine Grundform für den Modul q gewählt. Bezeichnen wir die aus den h ersten Reihen der Form B' und der Form φ' gebildeten symmetrischen Unterdeterminanten durch $\sigma_h' d_{h-1}' B_h'$ und durch $\tau_h' d_{h-1}' \varphi_h'$, so bestehen die Beziehungen $\sigma_h' B_h' = \tau_h' \varphi_h'$ ($h \leq n-2$) und $(-1)^I \cdot B_{n-1}' = \frac{b}{\sigma_1} = \delta \cdot m$. Die Größe $(-1)^I \cdot \tau_{n-2}' \varphi_{n-2}'$ stimmt mit dem Koeffizienten c^* der Form $\delta \cdot \varphi$ überein, so daß wir die Kongruenz (62) auch

$$(63) \quad -(-1)^I \cdot \sigma_1 \tau_{n-2}' \varphi_{n-2}' \equiv \tilde{b}_0^2 \pmod{\sigma_1 b}$$

schreiben können.

1. Bedeutet zunächst q eine ungerade Primzahl p , welche nicht in b und nicht in Δ aufgeht, so besitzt weder die Form φ' noch die Form B' einen Charakter C_p .

2. Zweitens sei q eine ungerade Primzahl p , welche in b aufgeht und mithin zu Δ relativ prim ist. Alsdaun besitzt die Form B' keine Charaktere C_p , während die Form φ' den einzigen Charakter $\left(\frac{\varphi_{n-2}'}{p}\right)$

liefert, welcher wegen der Kongruenz (63) den Wert $\left(\frac{-(-1)^I \cdot \sigma_1 \tau_{n-2}'}{p}\right)$ erhält.

3. Wenn drittens q eine ungerade Primzahl p ist, welche in Δ aufgeht und mithin zu b relativ prim ausfällt, so stellt die Form B' eine Grundform für den Modul p vor, sobald wir für φ' eine Grundform für diesen Modul genommen haben. Es besitzt die Form φ' einen Charakter $\left(\frac{\varphi_h'}{p}\right)$, wenn die Invariante e_h' durch p teilbar ist, und die Form B' einen Charakter $\left(\frac{B_h'}{p}\right)$, wenn die Invariante o_h' durch p teilbar ist. Nun ist

offenbar ein jedes e'_h ($h \leq n-2$) dann und nur dann durch p teilbar, wenn das zugehörige o'_h durch p teilbar wird. Demnach zieht ein jeder Charakter $\left(\frac{B'_h}{p}\right)$ einen bestimmten Charakter $\left(\frac{\varphi'_h}{p}\right)$ nach sich, und es sind diese beiden Charaktere wegen $\sigma'_h B'_h = \tau'_h \varphi'_h$ gegenseitig durcheinander bestimmt. Es muß ferner, wenn $o'_{n-1} \equiv 0 \pmod{p}$ wird, der Charakter $\left(\frac{B'_{n-1}}{p}\right)$ gleich $\left(\frac{(-1)^I \cdot \frac{b}{\sigma_1}}{p}\right)$ sein.

4. Endlich untersuchen wir die Charaktere für einen Modul $q^t = 2^t$.

I. Es möge zuerst $m \equiv 1 \pmod{2}$ sein. Gehört dann die Form φ' zu der Ordnung $O'_I(b)$, gelten sonach die Beziehungen $\sigma'_h = \tau'_h$ ($h \leq n-2$), so stimmen die Zahlen φ'_h mit den entsprechenden Größen B'_h überein, und es werden die Größen $\tau'_{h-1} e'_h \tau'_{h+1}$ ($h \leq n-2$) nur dann den Faktor 4 oder 8 enthalten, wenn die entsprechenden Größen $\sigma'_{h-1} o'_h \sigma'_{h+1}$ ($h \leq n-2$) durch 4 oder durch 8 teilbar sind. Infolgedessen wird jeder Charakter $C_4, C_8, \mathfrak{C}_4$ der Form φ' gleich einem bestimmten Charakter $C_4, C_8, \mathfrak{C}_4$ der Form B' sein. Außerdem wird, wie wir leicht bemerken, wenn

$\sigma'_{n-2} o'_{n-1} \equiv 0 \pmod{4}$ ist, der Charakter $(-1)^{\frac{B'_{n-1}-1}{2}}$ gleich $(-1)^{\frac{(-1)^I \cdot \frac{b}{\sigma_1} - 1}{2}}$, und wenn $\sigma'_{n-2} o'_{n-1} \equiv 0 \pmod{8}$ ist, der Charakter $\left(\frac{2}{B'_{n-1}}\right)$ gleich $\left(\frac{2}{\frac{b}{\sigma_1}}\right)$.

Ferner liefert, falls $\sigma_1 = 2$ ist, die Kongruenz (63) die Bedingung $-(-1)^I \cdot o_1 \varphi'_{n-2} \equiv 1 \pmod{4}$ oder

$$(-1)^{\frac{\varphi'_{n-2}-1}{2}} = (-1)^{\frac{-(-1)^I o_1 - 1}{2}}.$$

Wenn die Form φ' der Ordnung $O'_{II}(b)$ angehört, so können wir ähnlich schließen, oder wir können auch auf die folgende Art vorgehen: Im Falle einer Ordnung $O'_{II}(b)$ ist jedenfalls $\sigma_1 = 1$ und $b = \delta \cdot \sigma_1 m \equiv 1 \pmod{2}$. Infolgedessen können wir in diesem Falle die Zahlen $b_1, b_2, \dots, b_{n-1}$ sämtlich kongruent Null $\pmod{2^{\omega_1' + \omega_2' + \dots + \omega'_{n-1}} = 2^{\nu'_{n-1}}}$ annehmen. Alsdann werden wegen der Gleichungen

$$-b b' = \sum_{i=1}^{n-1} b_{n-i} b'_i - \prod_{h=1}^{n-1} o_h, \quad -b b'_k = \sum_{i=1}^{n-1} b_{n-i} b'_{ik}$$

auch die Zahlen $b', b'_1, b'_2, \dots, b'_{n-1}$ durch $2^{\nu'_{n-1}}$ teilbar sein. Daraus erhellt, daß für jede ganze Zahl z' die Kongruenz $B' \equiv z' \pmod{2^{\nu'_{n-1}}}$ genau $2^{\nu'_{n-1}}$ mal soviel Lösungen besitzen wird als die Kongruenz $\varphi' \equiv z' \pmod{2^{\nu'_{n-1}}}$. Nun können nach Kap. IX, wie wir wissen, die Charaktere der Form B' und die der Form φ' für die Moduln 2^t aus den Anzahlen

der Lösungen der verschiedenen Kongruenzen $B' \equiv z' \pmod{2^{n'-1}}$ und $\varphi' \equiv z' \pmod{2^{n'-1}}$ hergeleitet werden. Es sind daher in dem vorliegenden Falle wirklich die Charaktere der beiden Formen φ' und B' für die Moduln $2'$ gegenseitig durcheinander gegeben.

II. Es möge endlich $m \equiv 0, \Delta \equiv 1 \pmod{2}$ sein. Die Form φ' besitzt dann, falls $\sigma_1 b \equiv 0 \pmod{4}$ ist, einen Charakter $(-1)^{\frac{\varphi'_{n-2}-1}{2}}$, welcher infolge der Kongruenz (63) den Wert $(-1)^{\frac{-(-1)^f o_1 - 1}{2}}$ besitzt, und falls $\sigma_1 b \equiv 0 \pmod{8}$ ist, einen Charakter $\left(\frac{2}{\varphi'_{n-2}}\right)$, welcher wegen (63) gleich $\left(\frac{2}{o_1}\right)$ ist. Alle übrigen Charaktere der Form φ' für die Moduln $2'$ können durch diese Einheiten $(-1)^{\frac{\varphi'_{n-2}-1}{2}}$ und $\left(\frac{2}{\varphi'_{n-2}}\right)$ und durch Charaktere C_p ausgedrückt werden.

Kap. XIX. Über den Inbegriff der Darstellungen einer ganzen Zahl durch die verschiedenen Formen eines Genus G .

Bekanntlich gibt es nur eine endliche Anzahl von Formenklassen, welche dieselbe Determinante Δ und denselben Index I aufweisen. Da nun die Formen, welche einem und demselben Genus angehören, gewiß dieselbe Determinante und denselben Index besitzen, so wird infolgedessen auch ein jedes Genus nur eine endliche Anzahl verschiedener Klassen besitzen.

Es sei ein bestimmtes Genus G von Formen gegeben; wir greifen aus jeder seiner Klassen irgendeinen Repräsentanten heraus. Wenn das Genus G sich im ganzen aus g verschiedenen Formenklassen zusammensetzt, bekommen wir auf diese Weise g untereinander nicht-äquivalente Formen $f_1, f_2, \dots, f_g$, und es ist klar, daß eine jede Form des Genus G einer und nur einer dieser g Formen äquivalent ist. Daher wird das System dieser g Formen $f_1, f_2, \dots, f_g$ ein *vollständiges Formensystem* für das Genus G genannt. Man erkennt leicht, daß die den Formen $f_1, f_2, \dots, f_g$ adjungierten Formen $f'_1, f'_2, \dots, f'_g$ ein vollständiges Formensystem für das dem Genus G adjungierte Genus G' bilden.

Bezeichnen wir durch $\sigma_1 \varphi_1$ irgendeine durch Formen f des Genus G darstellbare ganze Zahl, für welche φ_1 zu Δ relativ prim ausfällt. Eine jede andere Zahl b , welche ebenfalls durch Formen f des Genus G darstellbar ist, muß alsdann einer bestimmten Kongruenz von der Form $b \equiv \sigma_1 \varphi_1 Z^2 \pmod{\sigma_0 o_1 \sigma_2}$ Genüge leisten. Es liege eine solche Zahl $b = \delta \cdot \sigma_1 m$ vor, für welche m zu Δ relativ prim ist, und wir wollen die Gesamtheit der primären Darstellungen dieser Zahl b durch die ver-

schiedenen Formen f eines vollständigen Formensystems des Genus G betrachten.

Erinnern wir uns, daß wir in Kap. XVIII für jede Zahl $b = \delta \cdot \sigma_1 m$ (m relativ prim zu Δ) ein bestimmtes Genus $G'_I(b)$, und falls $\sigma_1 = 1$, $m \equiv \kappa_1 \equiv 1 \pmod{2}$, $\kappa_1 > 1$ oder $\sigma_1 = 1$, $m \equiv \kappa_1 \equiv 0 \pmod{2}$, $\kappa_1 > 2$ ist, außerdem ein bestimmtes Genus $G'_{II}(b)$ definiert haben. Wir verschaffen uns jetzt ein vollständiges System von Formen φ' für das Genus $G'_I(b)$ und, falls das Genus $G'_{II}(b)$ zulässig ist, auch für das Genus $G'_{II}(b)$. Als dann wird jeder primären Darstellung der Zahl b durch eine der Formen f eine einzige Gruppe von Darstellungen einer einzigen dieser Formen φ' durch eine der Formen f' adjungiert sein. Wir erhalten mithin sämtliche möglichen primären Darstellungen der Zahl b durch die f , indem wir für jede der Formen φ' die sämtlichen nicht-äquivalenten Gruppen von Darstellungen durch die Formen f' aufsuchen. Für eine bestimmte Form $\varphi' = \{b'_{ik}\}$ kann dies auf folgende Weise geschehen:

Wir denken uns der Einfachheit halber φ' so angenommen, daß die ihr adjungierte Form $\varphi = \delta \cdot \{c_{ik}\}$ einen Hauptrest für den Modul b vorstellt, und bezeichnen durch $(-1)^I \cdot \delta \tau'_{n-2} \varphi'_{n-2}$ den ersten Koeffizienten dieser Form φ . Infolge der besonderen Wahl des Genus $G'_I(b)$ oder zutreffendenfalls des Genus $G'_{II}(b)$ ist gewiß die Kongruenz

$$-(-1)^I \cdot o_1 \tau'_{n-2} \varphi'_{n-2} \equiv \tilde{b}_0^2 \pmod{\sigma_1 b}$$

auflösbar, und wir schließen hieraus mittels der in Kap. XVIII gegebenen Sätze sofort, daß auch die $\frac{n(n-1)}{2}$ Kongruenzen

$$(50) \quad -o_1 c_{ik} \equiv b_i b_k \pmod{b}$$

lösbar sein werden. Es sei N die Anzahl der sämtlichen inkongruent Lösungen $(b_1, b_2, \dots, b_{n-1}) \pmod{b}$ dieser Kongruenzen. Enthält der Modul b genau μ verschiedene ungerade Primzahlen, so ist die Zahl N , wie wir sahen, gleich 2^μ , wenn $b \equiv 1 \pmod{2}$ oder $b \equiv 2 \pmod{4}$ ist, gleich $2^{\mu+1}$, wenn $b \equiv 4 \pmod{8}$ ist, und gleich $2^{\mu+2}$, wenn $b \equiv 0 \pmod{8}$ ist. Eine jede Gruppe von eigentlichen Darstellungen der Form φ' durch eine der Formen $f'_1, f'_2, \dots, f'_y$ muß jetzt zu einem und nur zu einem dieser N inkongruenten Lösungssysteme $(b_1, b_2, \dots, b_{n-1}) \pmod{b}$ der Kongruenzen (50) gehören. Um die Darstellungen von φ' zu finden, welche zu einer bestimmten dieser N Wurzeln $(b_1, b_2, \dots, b_{n-1})$ gehören, können wir folgendermaßen vorgehen:

Wir setzen

$$\frac{o_1 c_{ik} + b_i b_k}{b} = b_{ik}$$

und

$$B = b \xi^2 + 2 \sum_{i=1}^{n-1} b_i \xi \xi_i + \sum_{i,k=1}^{n-1} b_{ik} \xi_i \xi_k.$$

Da wir gezeigt haben, daß der Index, die Invarianten und die Charaktere des Genus G' auf eindeutig bestimmte Weise durch den Index, die Invarianten und die Charaktere des Genus $G'_I(b)$ oder des Genus $G'_{II}(b)$ dargestellt werden können, muß jetzt das Genus der Form B mit dem gegebenen Genus G identisch sein. Daher besitzt die der Form B adjungierte Form die Gestalt

$$B' = \sum_{i,k=1}^{n-1} b'_{ik} \xi'_i \xi'_k + 2 \sum_{i=1}^{n-1} b'_i \xi'_i \xi'_2 + b' \xi'^2_2$$

und gehört dem Genus G' an. Diese Form B' wird daher einer und nur einer der g Formen $f'_1, f'_2, \dots, f'_g$ äquivalent sein, die wir mit f' bezeichnen wollen. Es gibt dann keine Darstellung von φ' durch eine der übrigen $g-1$ von f' verschiedenen Formen f'_k , welche zu der Wurzel (b_h) gehört. Dagegen sind wohl Darstellungen von φ' durch die Form f' vorhanden, welche zu der Wurzel (b_h) gehören, und es liefert eine jede Substitution

$$(t'): \quad x'_i = \sum_{k=1}^{n-1} \vartheta'_i{}^k \xi'_k + t'_i \xi'_2 \quad (i = 1, 2, \dots, n)$$

von der Determinante 1, welche die Form f' in B' verwandelt, eine derartige Darstellung

$$(\vartheta'): \quad x'_i = \sum_{k=1}^{n-1} \vartheta'_i{}^k \xi'_k \quad (i = 1, 2, \dots, n)$$

der Form φ' durch f' . Ferner ist aus Kap. XVII bekannt, daß zwei verschiedene Transformationen (t') von f' in B' stets zu zwei verschiedenen Darstellungen von φ' durch f' führen. Infolgedessen ist die Anzahl der sämtlichen verschiedenen Darstellungen von φ' durch f' , welche zu der gegebenen Wurzel (b_h) gehören, gleich der Anzahl der verschiedenen Substitutionen von der Determinante 1, welche die Form f' in B' , oder, was auf dasselbe hinauskommt, welche die Form f' in sich überführen.

Diese sämtlichen Darstellungen lassen sich dann in eine gewisse Anzahl R nicht-äquivalenter Darstellungsgruppen verteilen. Eine jede dieser R Gruppen enthält genau soviel verschiedene Darstellungen (ϑ') , als man verschiedene Substitutionen von der Determinante 1 bilden kann, durch welche die Form φ' in sich selbst übergeht. Den R verschiedenen Gruppen von Darstellungen der Form φ' durch f' sind endlich R verschiedene Darstellungen der Zahl b durch die Form f adjungiert, welche alle zu der Wurzel $(b_1, b_2, \dots, b_{n-1})$ gehören.

Kap. XX. Maß eines positiven Genus. — Maß der Darstellungen einer ganzen Zahl durch die Formen eines positiven Genus.

Wir wollen jetzt unsere weiteren Untersuchungen auf den Fall $I = 0$, d. i. auf den Fall positiver Formen f und f' beschränken. Die Anzahl der ganzzahligen Substitutionen von der Determinante 1, vermittels deren eine positive Form f in sich selbst übergeführt werden kann, ist, wie man weiß, stets eine endliche. Wir bezeichnen diese Zahl durch $t(f)$. Die Größe $t(f)$ nimmt, wie man weiß, für alle Formen f derselben Klasse den gleichen Wert an. Ferner erkennt man leicht, daß für zwei adjungierte Formen f und f' die Größen $t(f)$ und $t(f')$ gleich sind. In der Tat, wird die Form f durch eine Substitution S in sich selbst transformiert, so geht die Form f' nach den in Kap. XII aufgestellten Sätzen durch die adjungierte Substitution S' in sich selbst über.

Die Größe $\frac{1}{t(f)}$ heißt das *Maß* der Klasse f^*). Bedeuten $f_1, f_2, \dots, f_g$ verschiedene Formenklassen, so wird die Summe der Maße

$$\frac{1}{t(f_h)} \quad (h = 1, 2, \dots, g)$$

der einzelnen Formen f_h das Maß des Klassensystems $f_1, f_2, \dots, f_g$ genannt. Zum Beispiel ist das Maß eines Genus G die Summe der Maße aller in diesem Genus enthaltenen Klassen und das Maß einer Ordnung O die Summe der Maße aller in dieser Ordnung enthaltenen Klassen. Hiernach ist das Maß einer Ordnung O gleich der Summe der Maße der sämtlichen in dieser Ordnung enthaltenen Genera G .

Durch eine positive Form f können nur positive Zahlen b dargestellt werden, und durch eine positive Form f' können nur positive Formen φ' dargestellt werden.

Wenn eine Zahl b vermittels eines Systems $x_i = t_i$ durch eine Form f dargestellt wird, so wird die Größe $\frac{1}{t(f)}$ als das *Maß* dieser Darstellung bezeichnet. Wenn mehrere Darstellungen (t) einer oder mehrerer Zahlen b durch eine oder mehrere Formen f vorliegen, so wird die Summe der Maße aller dieser einzelnen Darstellungen (t) das Maß des ganzen vorliegenden Systems von Darstellungen genannt. Hat man beispielsweise R verschiedene Darstellungen derselben Zahl b durch dieselbe Form f , so wird $\frac{R}{t(f)}$ das Maß dieser Darstellungen sein.

Wir betrachten jetzt wieder den Inbegriff der sämtlichen primären Darstellungen einer Zahl $b = \delta \cdot \sigma_1 m [\equiv \sigma_1 \varphi_1 Z^2 \pmod{\sigma_0 \sigma_1 \sigma_2}]$, m relativ prim zu Δ] durch ein vollständiges System von Formen $f_1, f_2, \dots, f_g$ eines

*) Eisenstein, Crelles Journal, Bd 35. [[Math. Abhandlungen, S. 180.]]

Genus G . Aber wir wollen jetzt die Zahl b und das Genus G als positiv voraussetzen ($\delta = 1, I = 0$). Es sei wiederum φ' irgend eine Form aus dem Genus $G'_I(b)$ oder auch, falls $\sigma_1 = 1, m \equiv \kappa_1 \equiv 1 \pmod{2}, \kappa_1 > 1$ oder $\sigma_1 = 1, m \equiv \kappa_1 \equiv 0 \pmod{2}, \kappa_1 > 2$ ist, aus dem Genus $G'_{II}(b)$, und es bedeute $\varphi = \{c_{ik}\}$ die zu φ' adjungierte Form. Aus jeder Lösung der $\frac{n(n-1)}{2}$ Kongruenzen

$$(50) \quad -o_1 c_{ik} \equiv b_i b_k \pmod{b}$$

entspringen (nach Kap. XIX) $R = \frac{t(f')}{t(\varphi')}$ verschiedene Darstellungen der Zahl b durch eine bestimmte Form f aus der Reihe der g Formen $f_1, f_2, \dots, f_g$. Für das Maß M_0 dieser R Darstellungen erhält man daher

$$M_0 = \frac{t(f')}{t(\varphi')} \cdot \frac{1}{t(f)} = \frac{1}{t(\varphi')}.$$

Die Größe M_0 ist also gleich dem Maße der Form φ' und fällt demnach unabhängig von der besonderen Form f aus. Infolgedessen liefert jede der verschiedenen Wurzeln ($b_1, b_2, \dots, b_{n-1}$) der Kongruenzen (50) ein gleiches Maß von Darstellungen der Zahl b durch die Formen des Genus G . Es ist dieses eben derjenige Umstand, durch welchen die Einführung des Maßbegriffes der Formen und Darstellungen eine so hohe Bedeutung gewinnt. Das Maß aller Darstellungen von b , welche zu den N verschiedenen Wurzeln der Kongruenzen (50) gehören, wird nun offenbar gleich $\frac{N}{t(\varphi')}$, d. i. gleich dem N -fachen Maße der Form φ' . Demgemäß wird dann das Maß aller derjenigen Darstellungen von b , welche aus den verschiedenen Klassen φ' des Genus $G'_I(b)$ oder zutreffendenfalls des Genus $G'_{II}(b)$ entspringen, gleich dem N -fachen Maße aller Formenklassen des Genus $G'_I(b)$, resp. des Genus $G'_{II}(b)$, d. i. gleich dem N -fachen Maße des Genus $G'_I(b)$, resp. des Genus $G'_{II}(b)$. Wir gewinnen dadurch den folgenden Satz:

Das Maß aller primären Darstellungen einer Zahl $b = \sigma_1 m$ (m relativ prim zu Δ) durch die verschiedenen Formen f des Genus G ist im allgemeinen gleich dem N -fachen Maß des Genus $G'_I(b)$ und im besonderen, wenn $\sigma_1 = 1, m \equiv \kappa_1 \pmod{2}, \kappa_1 > 2$ wird, gleich dem N -fachen Maß der beiden Genera $G'_I(b)$ und $G'_{II}(b)$.

Kap. XXI. Über die Anzahl der Darstellungen einer ganzen Zahl durch eine Summe von fünf Quadraten.

Die Anwendung des zuletzt gewonnenen Resultates auf den Fall $n = 5, \Delta = 1$ verschafft uns einige interessante Sätze über die Darstellung von ganzen Zahlen durch eine Summe von fünf Quadraten.

Bekanntlich bilden die positiven Formen mit fünf Variablen von der Determinante 1 eine einzige Formenklasse, welche durch die Form

$$f = x_1^2 + x_2^2 + x_3^2 + x_4^2 + x_5^2$$

repräsentiert werden kann. Diese Form f gehört dem einzigen Geschlecht G der Ordnung

$$O: \left(\begin{array}{l} \sigma_1=1, \sigma_2=1, \sigma_3=1, \sigma_4=1 \\ \sigma_1=1, \sigma_2=1, \sigma_3=1, \sigma_4=1 \end{array} \right), I=0$$

an, und sie repräsentiert zugleich, da alle Formen dieses Genus die Determinante 1 besitzen müssen, die einzige Klasse dieses Genus.

Der Form f ist die Form

$$f' = x_1'^2 + x_2'^2 + x_3'^2 + x_4'^2 + x_5'^2$$

adjungiert, welche mit f identisch ist. Das Maß der Klassen f und f' oder des Genus G ist, wie man ohne weiteres erkennt, gleich

$$\frac{1}{1 \cdot 2 \cdot 3 \cdot 4 \cdot 5 \cdot 2^4} = \frac{1}{2^7 \cdot 3 \cdot 5} = \frac{1}{1920} = \frac{1}{M_5}.$$

Bezeichnen wir für einen Augenblick die Anzahl der sämtlichen Systeme x_i ohne gemeinsamen Teiler, welche einer Gleichung $f(x_i) = m$ genügen, mit $(m)_5$, so ist das Maß der eigentlichen Darstellungen einer Zahl m durch eine Form des Genus G gleich $\frac{(m)_5}{M_5}$. Da $\Delta = 1$ ist, so wird eine jede beliebige Zahl m zu Δ relativ prim, und es können die Größen $\frac{(m)_5}{M_5}$ nach dem in Kap. XX bewiesenen Satze durch das Maß des einen Genus $G'_I(m)$ oder der beiden Genera $G'_I(m)$ und $G'_{II}(m)$ von Formen mit vier Variablen ausgedrückt werden. Wir haben infolgedessen, um zu einer Kenntnis der Größen $(m)_5$ zu gelangen, nur die Maße dieser Genera aufzusuchen.

Die Form f ist ein Hauptrepräsentant für den Modul 2; es ist $\sigma_1 = 1$ und $\alpha_1 = 5 \equiv 1 \pmod{2}$. Wir müssen demnach für eine Zahl m die Darstellungen $x_i = t_i$, in welchen die fünf Zahlen t_i nicht alle ungerade sind, und die Darstellungen $x_i = t_i$, in welchen $t_1 \equiv t_2 \equiv t_3 \equiv t_4 \equiv t_5 \equiv 1 \pmod{2}$ ist, voneinander unterscheiden. Die Darstellungen (t) der ersten Art sind mit Darstellungen von primitiven Formen φ' des Genus

$$G'_I(m): \left(\begin{array}{l} \tau_1' = 1, \tau_2' = 1, \tau_3' = 1 \\ e_1' = 1, e_2' = 1, e_3' = m \end{array} \right), J' = 0; \quad -\varphi_3' \equiv X^2 \pmod{m}$$

adjungiert, während die Darstellungen (t) der zweiten Art mit Darstellungen primitiver Formen φ' des Genus

$$G'_{II}(m): \left(\begin{array}{l} \tau_1' = 2, \tau_2' = 1, \tau_3' = 2 \\ e_1' = 1, e_2' = 1, e_3' = m \end{array} \right), J' = 0; \quad -2\varphi_3' \equiv X^2 \pmod{m}$$

adjungiert sind. Bezeichnet N die Anzahl der Lösungen der Kongruenz

$X^2 \equiv 1 \pmod{m}$, so ist infolgedessen die Anzahl $(m)_5^1$ der Darstellungen erster Art gleich dem $M_5 \cdot N$ -fachen Maße des Genus $G'_I(m)$ und die Anzahl $(m)_5^2$ der Darstellungen zweiter Art gleich dem $M_5 \cdot N$ -fachen Maße des Genus $G'_{II}(m)$.

Wie man leicht erkennt, gibt es für eine Zahl m nur dann Darstellungen (t) der zweiten Art, wenn $m \equiv 5 \pmod{8}$ ist. Denn die Kongruenzen $t_i \equiv 1 \pmod{2}$ ($i = 1, 2, 3, 4, 5$) ergeben unmittelbar $t_i^2 \equiv 1 \pmod{8}$ und $m = \sum_{i=1}^5 t_i^2 \equiv 5 \pmod{8}$. — Die in Kap. XI aufgestellten Bedingungen

für die Existenz eines Genus lassen erkennen, daß das Genus $G'_I(m)$ für jede Zahl m und das Genus $G'_{II}(m)$ nur für ein $m \equiv 5 \pmod{8}$ Formen enthält. Es wird demnach eine jede Zahl m als eine Summe von fünf beliebigen Quadraten darstellbar sein, während nur die Zahlen $m \equiv 5 \pmod{8}$ sich als eine Summe von fünf ungeraden Quadraten darstellen lassen.

Alle Darstellungen einer Zahl m durch eine Summe von fünf Quadraten, bei welchen die fünf einzelnen Quadrate abgesehen von ihrer Reihenfolge und den Vorzeichen ihrer Wurzeln x_i miteinander übereinstimmen, fassen wir als eine einzige *Zerfällung* der Zahl m in eine Summe von fünf Quadraten zusammen. Bezeichnen wir die Anzahl der Zerfällungen einer Zahl m in eine Summe von fünf Quadraten durch D_5 , die Anzahl der sämtlichen Darstellungen von m durch eine Summe

$$n_1 x_1^2 + n_2 x_2^2 + \dots + n_\nu x_\nu^2 \quad (n_1 + n_2 + \dots + n_\nu \leq 5)$$

durch $R(n_1, n_2, \dots, n_\nu)$, so gilt die Relation

$$\begin{aligned} D_5 = & \frac{1}{3840} R(1, 1, 1, 1, 1) + \frac{1}{192} R(1, 1, 1, 2) + \frac{1}{768} R(1, 1, 1, 1) \\ & + \frac{1}{48} R(1, 1, 3) + \frac{1}{64} R(1, 2, 2) + \frac{1}{64} R(1, 1, 2) + \frac{1}{128} R(1, 1, 1) \\ & + \frac{1}{16} R(1, 4) + \frac{5}{48} R(2, 3) + \frac{1}{24} R(1, 3) + \frac{1}{64} R(2, 2) + \frac{3}{64} R(1, 2) \\ & + \frac{5}{128} R(1, 1) - \frac{1}{40} R(5) + \frac{1}{16} R(4) + \frac{1}{64} R(2) + \frac{35}{256} R(1). \end{aligned}$$

Kap. XXII. Über die Bestimmung des Maßes einiger positiver Genera.

Unter Anwendung der von Dirichlet*) gegebenen Prinzipien können wir das Maß eines beliebigen Genus G positiver Formen bestimmen. Ich wollte zeigen, worauf sich allgemein diese Bestimmung gründet; wegen der Kürze der Zeit muß ich mich jedoch darauf beschränken, in dem

*) Dirichlet, Recherches sur diverses applications de l'analyse infinitésimale à la théorie des nombres, Crelles Journal, Bd. 19 und 21. [[Werke, Bd. I.]]

Folgenden nur die Maße einiger spezieller Genera mit drei und vier Variablen zu betrachten.

I. Wir gehen dabei aus von Summen der Form

$$S = \sum_{x_i} \frac{1}{\{f(x_i)\}^{\frac{n}{2}(1+\varrho)}}, \quad (n \geq 2)$$

in welchen $f(x_i)$ eine beliebige positive Form von einer Determinante $\Delta > 0$ und ϱ eine beliebige positive GröÙe bedeutet, und in welchen die Variablen $x_1, x_2, \dots, x_n$ Systeme von n ganzen, nicht sämtlich verschwindenden Zahlen durchlaufen sollen.

1. Es sei N eine ganze positive Zahl und $\alpha_1, \alpha_2, \dots, \alpha_n$ irgendwelche n Reste für diesen Modul N . Wir wollen zunächst für die GröÙen $x_1, x_2, \dots, x_n$ alle Systeme von ganzen Zahlen einsetzen, welche nach dem Modul N die Reste $\alpha_1, \alpha_2, \dots, \alpha_n$ lassen, das sind alle Systeme von der Gestalt

$$(x) \quad x_i = Nv_i + \alpha_i, \quad (v_i = -\infty, \dots, -1, 0, 1, \dots, +\infty)$$

(wobei das System $x_1 = 0, x_2 = 0, \dots, x_n = 0$, falls es gleichfalls von der Form $Nv_i + \alpha_i$ ($i = 1, 2, \dots, n$) ist, stets ausgeschlossen wird). Da die Form f positiv ist, so nimmt für jedes einzelne dieser Wertesysteme (x)

der Ausdruck $\{f(x_1, x_2, \dots, x_n)\}^{\frac{n}{2}}$ einen positiven und von Null verschiedenen Wert an. Wir wollen die Anzahl aller derjenigen unter diesen Systemen, für welche die GröÙe $\{f(x_i)\}^{\frac{n}{2}}$ nicht größer ausfällt als eine positive GröÙe t , durch τ bezeichnen. Anstelle der Ungleichung

$$\left(\sum_{i,k=1}^n a_{ik} x_i x_k \right)^{\frac{n}{2}} \leq t$$

können wir schreiben

$$\sum_{i,k=1}^n a_{ik} \frac{x_i}{t^{\frac{1}{n}}} \frac{x_k}{t^{\frac{1}{n}}} \leq 1,$$

oder, indem wir

$$\xi_i = \frac{x_i}{N t^{\frac{1}{n}}} = \frac{1}{t^{\frac{1}{n}}} v_i + \frac{\alpha_i}{N t^{\frac{1}{n}}}; \quad a_{ik} N^2 = A_{ik}$$

setzen, auch

$$\sum_{i,k=1}^n A_{ik} \xi_i \xi_k \leq 1.$$

Der Grenzwert des Verhältnisses $\frac{\tau}{t}$ für ein unendlich wachsendes t wird nach einem bekannten Satze von Dirichlet gleich dem n -fachen Integral

$$J = \int \int \dots \int d\xi_1 d\xi_2 \dots d\xi_n,$$

in welchem die Variablen ξ_i alle diejenigen Wertsysteme zu durchlaufen haben, für welche die Ungleichung

$$\sum_{i,k=1}^n A_{ik} \xi_i \xi_k \leq 1$$

erfüllt ist. Dieses Integral erhält nach Dirichlet den Wert

$$J = \frac{1}{|A_{ik}|^{\frac{1}{2}}} \cdot \frac{\left\{ \Gamma\left(\frac{1}{2}\right) \right\}^n}{\Gamma\left(1 + \frac{n}{2}\right)},$$

in welchem $|A_{ik}|$ die aus den n^2 Größen A_{ik} gebildete Determinante bedeutet. Diese Determinante ist wegen $A_{ik} = a_{ik} N^2$ gleich $|a_{ik}| N^{2n} = \Delta \cdot N^{2n}$.

Ferner wird bekanntlich die Größe $\Gamma\left(\frac{1}{2}\right) = \pi^{\frac{1}{2}}$, während $\Gamma\left(1 + \frac{n}{2}\right)$, je nachdem n gerade oder ungerade ist, den Wert

$$1 \cdot 2 \cdot 3 \cdots \frac{n}{2} = \frac{n}{2} \cdot \frac{n-2}{2} \cdots \frac{4}{2} \cdot \frac{2}{2}$$

oder den Wert

$$\frac{1}{2} \cdot \frac{3}{2} \cdot \frac{5}{2} \cdots \frac{n}{2} \Gamma\left(\frac{1}{2}\right) = \pi^{\frac{1}{2}} \cdot \frac{n}{2} \cdot \frac{n-2}{2} \cdots \frac{3}{2} \cdot \frac{1}{2}$$

annimmt. Wir können mithin

$$\frac{\left\{ \Gamma\left(\frac{1}{2}\right) \right\}^n}{\Gamma\left(1 + \frac{n}{2}\right)} = \pi^{\left[\frac{n}{2}\right]} : \prod_{h=0}^{\left[\frac{n-1}{2}\right]} \left(\frac{n-2h}{2}\right) = e_n$$

setzen, und wir bekommen

$$J = e_n \cdot \frac{\Delta^{-\frac{1}{2}}}{N^n}.$$

Der Grenzwert des Quotienten $\frac{\tau}{t}$ für ein unendlich abnehmendes $\frac{1}{t^n}$ ist hiernach endlich.

Infolge dieses Umstandes konvergiert nach einem weiteren Satze von Dirichlet die unendliche Reihe

$$S = \sum_{[f(x_i)]^{\frac{n}{2}(1+\varrho)}} \frac{1}{\left(x_i = N v_i + \alpha_i, v_i = -\infty, \dots, -1, 0, 1, \dots, +\infty \right)} \left(x_1, x_2, \dots, x_n \neq (0, 0, \dots, 0) \right)$$

für ein jedes positive ϱ , und es wird der Grenzwert von $\varrho \cdot S$ für ein positives unendlich abnehmendes ϱ gleich

$$\lim_{\varrho=0} (\varrho \cdot S) = J = e_n \cdot \frac{\Delta^{-\frac{1}{2}}}{N^n}.$$

Ist der Modul N gleich 1, so haben die Variablen x_i alle möglichen Systeme von ganzen Zahlen mit Ausnahme des einen $x_1 = 0, x_2 = 0, \dots, x_n = 0$

zu durchlaufen, und der Limes von $\varrho \cdot S$ nimmt seinen größten Wert $e_n : \sqrt{\Delta}$ an.

2. Wir wollen jetzt annehmen, daß der größte gemeinsame Teiler der n Zahlen α_i zu N relativ prim wird, und wir wollen den n Größen v_i nur solche ganzzahligen Werte erteilen, für welche die n Zahlen $x_i = Nv_i + \alpha_i$ ohne gemeinsamen Teiler ausfallen. Anstatt der Summe S erhalten wir dann eine kleinere Summe S_0 .

Bezeichnen wir die verschiedenen in N aufgehenden Primzahlen durch $q_1, q_2, \dots, q_M$ und setzen wir

$$S_n = \sum_{z=1}^{\infty} \frac{1}{z^n} = \frac{1}{1^n} + \frac{1}{2^n} + \frac{1}{3^n} + \dots,$$

$$\varphi_n(N) = N^n \left(1 - \frac{1}{q_1^n}\right) \left(1 - \frac{1}{q_2^n}\right) \dots \left(1 - \frac{1}{q_M^n}\right) = N^n(N)_n,$$

so wird der Grenzwert von $\varrho \cdot S_0$ für unendlich abnehmendes ϱ gleich

$$\lim_{\varrho=0} (\varrho \cdot S_0) = \frac{e_n}{S_n} \cdot \frac{\Delta^{-\frac{1}{2}}}{\varphi_n(N)}.$$

Wir bemerken, daß wir den hier auftretenden Ausdruck $\varphi_n(N)$ in ähnlicher Weise definieren können wie die speziellere Funktion

$$\varphi_1(N) = N \left(1 - \frac{1}{q_1}\right) \left(1 - \frac{1}{q_2}\right) \dots \left(1 - \frac{1}{q_M}\right).$$

Denn wir haben den Satz:

Lassen wir jede der n Zahlen α_i die sämtlichen N Werte $1, 2, \dots, N$ durchlaufen, so gibt es unter den N^n sich dabei ergebenden Systemen $\varphi_n(N)$ Systeme, für welche der größte Teiler der n Größen $\alpha_1, \alpha_2, \dots, \alpha_n$ zu N relativ prim wird.

Man kann leicht die folgenden Relationen beweisen, in denen die Größe m alle ganzen positiven und zu N relativ primen Zahlen durchlaufen soll:

$$1 = \lim_{\varrho=0} (\varrho S_{1+\varrho}) = \frac{N}{\varphi_1(N)} \lim_{\varrho=0} \left(\varrho \sum_m \frac{1}{m^{1+\varrho}} \right); \quad S_n = \frac{N^n}{\varphi_n(N)} \cdot \sum_m \frac{1}{m^n} \quad (n > 1).$$

3. Wir leiten jetzt eine Umformung her, welche für die weiteren Untersuchungen wichtig ist. Wir wollen mit m oder m_0 alle positiven und zu N relativ primen ganzen Zahlen und mit $q', q'', \dots, q^{(\mu)}$ die verschiedenen in einer Zahl m aufgehenden Primzahlen bezeichnen. Es seien $\psi(m)$ und $\Psi(m)$ zwei Funktionen, welche den Bedingungen

$$\psi(m \cdot m_0) = \psi(m) \cdot \psi(m_0); \quad \Psi(m \cdot m_0) = \Psi(m) \cdot \Psi(m_0),$$

$$\psi(1) = 1, \quad -1 \leq \psi(q) \leq 1; \quad \Psi(1) = 1, \quad \Psi(q) = \pm \frac{1}{q^{1+\varrho}}$$

genügen.

Die über alle verschiedenen Zahlen m ausgedehnte Summe

$$\sum_m = \sum_m \left\{ \prod_{k=1}^{\nu} (1 + \psi(q^{e_k}) \cdot \Psi(m_k)) \right\}$$

ist für jedes positive q konvergent.

Indem wir die Zahlen m in ihre verschiedenen Primzahlpotenzen q^i zerlegen, erkennen wir, daß die Summe $\sum$ gleich dem über sämtliche in N nicht aufgehenden Primzahlen q ausgedehnten Produkt

$$\prod_q [1 + (1 + \psi(q))\Psi(q) + (1 + \psi(q))\Psi(q^2) + (1 + \psi(q))\Psi(q^3) + \dots]$$

wird. Die Einzelglieder dieses Produktes sind wegen $\Psi(q^i) = [\Psi(q)]^i$ gleich

$$1 + \frac{(1 + \psi(q))\Psi(q)}{1 - \Psi(q)} = \frac{1 - [\psi(q) \cdot \Psi(q)]^2}{[1 - \Psi(q)][1 - \psi(q) \cdot \Psi(q)]},$$

sodaß wir

$$\sum = \frac{\prod_q (1 - \psi(q)) \prod_q (1 - \psi(q) \cdot \Psi(q))}{\prod_q (1 - [\psi(q) \cdot \Psi(q)]^2)}$$

bekommen. Nun gilt, wenn man $\Phi = \Psi$ oder $= \psi \cdot \Psi$ oder $= \psi^2 \cdot \Psi^2$ nimmt, stets

$$\prod_q \frac{1}{1 - \Phi(q)} = \sum_m \Phi(m),$$

und hieraus geht sofort die Identität

$$(64) \quad \sum = \frac{\sum_m \Psi(m) \cdot \sum_m \psi(m) \Psi(m)}{\sum_m (\psi(m) \Psi(m))^2}$$

hervor.

II. ($\alpha = 2$.) Bekanntlich haben zwei positive binäre Formen f , welche demselben Genus angehören, stets dasselbe Maß $\frac{1}{t(f)} \left(= \frac{1}{2} \text{ oder } = \frac{1}{4} \text{ oder } = \frac{1}{6} \right)$. Infolgedessen ist das Maß eines Genus f von Formen mit zwei Variablen gleich der mit der Konstanten $\frac{1}{t(f)}$ multiplizierten Klassenanzahl dieses Genus und kann daher mit Hilfe der von Dirichlet aufgestellten Formeln bestimmt werden. So findet sich z. B., wenn N irgendeine Zahl kongruent 1 (mod 4) ist und $P_1, P_2, \dots, P_\nu$ die verschiedenen in N aufgehenden Primzahlen sind, das Maß eines beliebigen Genus

$$\left(\begin{matrix} \sigma_1 = 1 \\ \sigma_1 = N \end{matrix} \right), \quad \left(\begin{matrix} \varphi \\ P_k \end{matrix} \right) \quad (k = 1, 2, \dots, \nu)$$

gleich

$$\frac{1}{2^\nu} \sqrt{N \cdot h(N)},$$

wo

$$h(N) = \frac{1}{\pi} \sum_m \left(\frac{-N}{m} \right) \frac{1}{m} \quad (m \text{ rel. pr. zu } 2N)$$

ist $[h(N) > 0]$. $(n = 3)$. Wir betrachten jetzt eine Ordnung

$$O: \begin{pmatrix} \sigma, & \sigma' \\ o, & o' \end{pmatrix}$$

von primitiven Formen mit drei Variablen. Dieser Ordnung wird die Ordnung

$$O': \begin{pmatrix} \sigma', & \sigma \\ o', & o \end{pmatrix}$$

adjungiert sein.

Wir beschäftigen uns insbesondere mit dem Fall, in welchem die Invarianten o und o' alle beide ungerade sind. In diesem Falle wird $\sigma = 1$, $\sigma' = 1$; $o \equiv 1$, $o' \equiv 1 \pmod{2}$. Wir bezeichnen die Primzahlen, welche in o aufgehen, durch $t_1, t_2, \dots, t_\delta$, die Primzahlen, welche in o' aufgehen, durch $t'_1, t'_2, \dots, t'_{\delta'}$, die Primzahlen, welche in o , aber nicht in o' aufgehen, durch $p_1, p_2, \dots, p_\delta$, die Primzahlen, welche in o' , aber nicht in o aufgehen, durch $p'_1, p'_2, \dots, p'_{\delta'}$, endlich die Primzahlen, welche sowohl in o als in o' aufgehen, durch $r_1 = r'_1, r_2 = r'_2, \dots, r_\tau = r'_\tau (\tau = \tau')$. Die Zahlen t bestehen aus den Zahlen p und den Zahlen r , die Zahlen t' aus den Zahlen p' und den Zahlen r' ; mithin ist $\delta = \delta + \tau$, $\delta' = \delta' + \tau'$.

1. Es möge Φ eine in der Ordnung O auftretende Grundform für den Modul $2oo'$ und Φ' ihre in der Ordnung O' auftretende adjungierte Form sein. Setzen wir den ersten Koeffizienten von Φ gleich φ und den ersten Koeffizienten von Φ' gleich φ' , so besitzen die Formen Φ und Φ' die $\vartheta + \vartheta'$ Charaktere

$$C(\varphi): \left(\frac{\varphi}{p_1} \right), \left(\frac{\varphi}{p_2} \right), \dots, \left(\frac{\varphi}{p_\delta} \right); \left(\frac{\varphi}{r_1} \right), \left(\frac{\varphi}{r_2} \right), \dots, \left(\frac{\varphi}{r_\tau} \right),$$

$$C'(\varphi'): \left(\frac{\varphi'}{p'_1} \right), \left(\frac{\varphi'}{p'_2} \right), \dots, \left(\frac{\varphi'}{p'_{\delta'}} \right); \left(\frac{\varphi'}{r'_1} \right), \left(\frac{\varphi'}{r'_2} \right), \dots, \left(\frac{\varphi'}{r'_{\tau'}} \right),$$

welche der Gleichung

$$\left(\frac{\varphi}{o} \right) \cdot \left(\frac{\varphi'}{o'} \right) = (-1)^{\frac{\varphi-1}{2} \cdot \frac{\varphi'-1}{2} + \frac{o+1}{2} \cdot \frac{\varphi-1}{2} + \frac{o'+1}{2} \cdot \frac{\varphi'-1}{2}}$$

oder

$$(-1)^{\frac{o'\varphi+1}{2} \cdot \frac{o\varphi'+1}{2}} = (-1)^{\frac{o+1}{2} \cdot \frac{o'+1}{2}} \left(\frac{\varphi}{o} \right) \cdot \left(\frac{\varphi'}{o'} \right) = \tau$$

genügen.

Wenn die Form Φ Hauptrepräsentant für einen Modul N ist, so läßt sie sich schreiben

$$\Phi \equiv \begin{pmatrix} \varphi, & 0, & 0 \\ 0, & \frac{o\varphi'}{\varphi}, & 0 \\ 0, & 0, & \frac{o o'}{\varphi'} \end{pmatrix} \pmod{N}.$$

Indem wir $N=t$ oder $N=p'$ annehmen, erkennen wir hieraus leicht, daß die Anzahl der Lösungen der Kongruenz

$$\Phi(\xi_1, \xi_2, \xi_3) \equiv 0 \pmod{t} \quad \text{resp.} \quad \Phi(\xi_1, \xi_2, \xi_3) \equiv 0 \pmod{p'}$$

gleich

$$t^2 \quad \text{resp.} \quad p'^2 + \left(\frac{-o\varphi'}{p'}\right) p'(p'-1)$$

ist. Dementsprechend wird die Anzahl der nach einem Modul t oder p' inkongruenten Systeme (ξ_1, ξ_2, ξ_3) , für welche die Form $\Phi(\xi_1, \xi_2, \xi_3)$ zu t resp. p' relativ prim ausfällt, gleich

$$t^3 \left(1 - \frac{1}{t}\right) \quad \text{resp.} \quad p'^3 \left(1 - \frac{1}{p'}\right) \left[1 - \left(\frac{-o\varphi'}{p'}\right) \frac{1}{p'}\right].$$

Ist der Modul N gleich 4, so hat die Anzahl der Lösungen $(\xi_1, \xi_2, \xi_3) \pmod{4}$ der Kongruenz

$$\Phi(\xi_1, \xi_2, \xi_3) \equiv o' \pmod{4}$$

den Wert $2^3(2-1)$ oder $2^3(2+1)$, je nachdem die Einheit

$$T = (-1)^{\frac{o\varphi'+1}{2} \cdot \frac{o'\varphi+1}{2}}$$

gleich 1 oder gleich -1 ist. Unter diesen $2^3(2-T)$ Lösungen der Kongruenz $\Phi(\xi_i) \equiv o' \pmod{4}$ befinden sich, wie man leicht erkennt, jedenfalls nicht die Systeme $\xi_1 \equiv 1, \xi_2 \equiv 1, \xi_3 \equiv 1 \pmod{2}$. Denn für diese Systeme wird stets

$$\Phi(\xi_i) \equiv \varphi + \frac{o\varphi'}{\varphi} + \frac{o o'}{\varphi'} \equiv -o' \pmod{4}.$$

2. Wir gehen jetzt zur Bestimmung des Maßes M eines Genus

$$G: \begin{pmatrix} 1, & 1 \\ o, & o' \end{pmatrix}, C(\varphi), C'(\varphi')$$

über, in welchem die Charaktere $C(\varphi), C'(\varphi')$ gegebene Werte (1 oder -1) besitzen. Wir bilden zunächst ein vollständiges Formensystem $\Phi_1, \Phi_2, \dots, \Phi_K$ für das Genus G .

Eine zu $2oo'$ relativ prime, positive Zahl m , welche $\equiv o' \pmod{4}$ ist, kann offenbar nur dann durch die Formen Φ dargestellt werden, wenn die Einheiten $C(m)$ gleich den gegebenen Charakteren $C(\varphi)$ sind. Sobald aber die sämtlichen Gleichungen $C(m) = C(\varphi)$ statthaben, wird, wenn wir annehmen, daß die Zahl m die μ ungeraden Primzahlen $q_1, q_2, \dots, q_\mu$ enthält, das Maß der eigentlichen Darstellungen von m durch die Formen Φ

gleich dem 2^μ -fachen Maße des Genus

$$\chi': \left(\frac{1}{o'm} \right), \left(\frac{\chi'}{t'} \right) = \left(\frac{\varphi'}{t'} \right), \left(\frac{\chi'}{q} \right) = \left(\frac{-o}{q} \right),$$

d. i. gleich

$$\frac{1}{2^{g'}} \sqrt{o'm} \cdot h(o'm)$$

sein.

Wir bilden jetzt die Summe

$$\sum_{\varrho} = \sum_{k=1}^K \left(\sum \frac{\{t(\Phi_k)\}^{-1}}{\{\Phi_k(\xi_1, \xi_2, \xi_3)\}^{\frac{3}{2}(1+\varrho)}} \right),$$

in welcher die Größen ξ_1, ξ_2, ξ_3 alle verschiedenen Systeme von ganzen Zahlen ohne gemeinsamen Teiler durchlaufen sollen, für welche $\Phi_k(\xi_1, \xi_2, \xi_3)$ zu $2oo'$ relativ prim und $\equiv o' \pmod{4}$ ausfällt. Diese Systeme ξ_1, ξ_2, ξ_3 sind, wie sich aus dem in 1. über die Kongruenzen $\Phi \equiv 0 \pmod{t}$ oder $\pmod{p'}$ und $\Phi \equiv o' \pmod{4}$ Bemerkten ergibt, in

$$\frac{1}{2} \left(1 - \frac{\tau}{2} \right) (4oo')^3 (2oo')_1 \cdot \prod_{p'} \left[1 - \left(\frac{-o\varphi'}{p'} \right) \frac{1}{p'} \right]$$

arithmetischen Reihen

$$\xi_1 = 4oo' V_1 + \Xi_1, \quad \xi_2 = 4oo' V_2 + \Xi_2, \quad \xi_3 = 4oo' V_3 + \Xi_3$$

enthalten. Der Grenzwert von $\varrho \cdot \sum_{\varrho}$ für ein unendlich abnehmendes ϱ wird infolgedessen gleich

$$L = \frac{\frac{1}{2} \left(1 - \frac{\tau}{2} \right) (4oo')^3 \cdot (2oo')_1 \cdot \prod_{p'} \left[1 - \left(\frac{-o\varphi'}{p'} \right) \frac{1}{p'} \right] \cdot e_3 \cdot M}{\sqrt{o^3 o'} (4oo')^3 (2oo')_3 \cdot S_3}.$$

Jetzt können wir die Summe $\sum_{\varrho}$ nach den numerischen Werten der Zahlen $\Phi_k(\xi_1, \xi_2, \xi_3)$ ordnen, und wir bekommen dadurch einen Ausdruck

$$\sum_{\varrho} = \sum \frac{\{m\}}{m^{\frac{3}{2}(1+\varrho)}}, \quad \begin{matrix} (m \text{ rel. pr. zu } 2oo'; \\ o'm \equiv 1 \pmod{4}) \end{matrix}$$

in welchem die Funktionen $\{m\}$ die Maße der Darstellungen der Zahlen m durch die sämtlichen Formen Φ bedeuten werden. Mithin können wir schreiben

$$\sum_{\varrho} = \frac{1}{2^{g'}} \cdot \sum \frac{\sqrt{o'm} \cdot h(o'm)}{m^{\frac{3}{2}(1+\varrho)}} \quad \begin{matrix} (C(m) = C(\varphi) \\ o'm \equiv 1 \pmod{4}) \end{matrix}$$

und

$$L = \lim_{\varrho=0} \left(\varrho \sum_{\varrho} \right) = \frac{1}{2^{g'}} \cdot \lim \left(\varrho \sum \frac{\sqrt{o'm} \cdot h(o'm)}{m^{\frac{3}{2}(1+\varrho)}} \right) \cdot \begin{matrix} (C(m) = C(\varphi) \\ o'm \equiv 1 \pmod{4}) \end{matrix}$$

Ein Vergleich der beiden für den Grenzwert L gewonnenen Ausdrücke ergibt die Formel

$$\begin{aligned}
 & \frac{\left(1 - \frac{1}{2}\right) (2oo')_1 \cdot \prod_{p'} \left[1 - \left(\frac{-o\varphi'}{p'}\right) \frac{1}{p'}\right] \cdot \prod_p \left[1 - \left(\frac{-o'\varphi}{p}\right) \frac{1}{p}\right] \cdot M \cdot e_3}{2oo' \cdot (2oo')_3 \cdot S_3} = \\
 & \frac{1}{2^{\mathfrak{J}'}} \cdot \lim \left(\varrho \sum \frac{\sqrt{m} h(o'm) \prod_p \left[1 - \left(\frac{-o'm}{p}\right) \frac{1}{p}\right]}{m^{\frac{3}{2}(1+\varrho)}} \right) = \\
 & \frac{1}{2^{\mathfrak{J}'}} \cdot \lim \left(\varrho \sum \frac{\sqrt{m} h(o^2o'm)}{m^{\frac{3}{2}(1+\varrho)}} \right). \\
 & [o^2o'm \equiv 1 \pmod{4}; \quad C(m) = C(\varphi)]
 \end{aligned}$$

Dieselbe verwandelt sich, indem wir

$$\begin{aligned}
 M &= \frac{oo'}{2^{\mathfrak{J}+\mathfrak{J}'}} \cdot \prod_{p'} \left[1 + \left(\frac{-o\varphi'}{p'}\right) \frac{1}{p'}\right] \cdot \prod_p \left[1 + \left(\frac{-o'\varphi}{p}\right) \frac{1}{p}\right] \cdot \\
 & \cdot \prod_{r=r'} \left(1 - \frac{1}{rr'}\right) \cdot \left(1 + \frac{1}{2}\right) \cdot M_0
 \end{aligned}$$

setzen, in

$$(65) \quad \frac{(2oo')_1 \cdot (2oo')_2 \cdot e_3 M_0}{2^{\mathfrak{J}} \cdot (2oo')_3 \cdot 2S_3} = \lim \left(\varrho \sum \frac{\sqrt{m} h(o^2o'm)}{m^{\frac{3}{2}(1+\varrho)}} \right) \cdot \left(\begin{matrix} o^2o'm \equiv 1 \pmod{4} \\ C(m) = C(\varphi) \end{matrix} \right)$$

Beachten wir nun, daß das Maß des Genus

$$G: \begin{pmatrix} 1, & 1 \\ o, & o' \end{pmatrix}, \quad C(\varphi), \quad C'(\varphi')$$

mit dem Maße des Genus

$$G': \begin{pmatrix} 1, & 1 \\ o', & o \end{pmatrix}, \quad C'(\varphi'), \quad C(\varphi)$$

übereinstimmt, so können wir sofort eine zweite Gleichung hinschreiben:

$$(65') \quad \frac{(2oo')_1 \cdot (2oo')_2 \cdot e_3 M_0}{2^{\mathfrak{J}'} \cdot (2oo')_3 \cdot 2S_3} = \lim \left(\varrho \sum \frac{\sqrt{m'} h(o^2o'm)}{m'^{\frac{3}{2}(1+\varrho)}} \right) \cdot \left(\begin{matrix} o'^2om \equiv 1 \pmod{4} \\ C'(m') = C'(\varphi') \end{matrix} \right)$$

Wir ersehen jetzt aus der Gleichung (65), daß die Größe M_0 von den speziellen Charakteren $C'(\varphi')$ unabhängig ist, und aus der Gleichung (65'), daß die Größe M_0 von den speziellen Charakteren $C(\varphi)$ unabhängig ist. Die Größe M_0 kann daher nur von den beiden Invarianten o und o' abhängen.

3. Um den Wert von M_0 zu finden, können wir auf folgende Weise verfahren:

Bilden wir die Summe der Gleichungen (65) für die sämtlichen Genera der Ordnung O , welche die nämlichen Charaktere $C'(\varphi')$ besitzen, so ergibt sich die Gleichung

$$\frac{(2oo')_1 \cdot (2oo')_2 \cdot e_3 M_0}{(2oo')_3 \cdot 2S_3} = \lim \left(\varrho \sum \frac{\sqrt{m} h(o^2o'm)}{m^{\frac{3}{2}(1+\varrho)}} \right) = \{o^2o'\}, \quad \left(\begin{matrix} o^2o'm \equiv 1 \pmod{4} \\ m \text{ rel. pr. zu } o^2o' \end{matrix} \right)$$

worin die Summation jetzt über alle zu $2o^2o'$ relativ primen Zahlen m , welche $\equiv o^2o' \pmod{4}$ sind, auszudehnen ist. Ebenso erhalten wir aus der Gleichung (65') die Relation

$$\frac{(2oo')_1 \cdot (2oo')_2 \cdot e_3 M_0}{(2oo')_3 \cdot 2S_3} = \lim \left(\varrho \sum \frac{\sqrt{m'} h(o'^2 o m')}{m'^{\frac{3}{2}(1+\varrho)}} \right) = \{o'^2 o\}, \quad \left(\begin{matrix} o'^2 o m' \equiv 1 \pmod{4} \\ m' \text{ rel. pr. zu } o'^2 o \end{matrix} \right)$$

worin die Summation über alle zu $2o'^2o$ relativ primen Zahlen m' , welche $\equiv o'^2o \pmod{4}$ sind, auszudehnen ist. Es ergibt sich demnach die Gleichung

$$\{o^2o'\} = \{o'^2o\}.$$

In dem Falle, daß eine der Invarianten o, o' gleich 1 ist, gewinnen wir daraus insbesondere die Formel

$$\{D^2\} = \{D\}, \quad [D \equiv 1 \pmod{2}]$$

mit deren Hilfe wir dann noch

$$\begin{aligned} \{o^2o'\} &= \{o^4o'^2\} = \{o^2o'^2\} = \{oo'\}, \quad \{o'^2o\} = \{o'^4o^2\} = \{o'^2o^2\} = \{o'o\}; \\ \{o^2o'\} &= \{o'^2o\} = \{oo'\} \end{aligned}$$

bekommen.

Wir wollen jetzt die Größe

$$\{D\} = \lim \left(\varrho \sum_R \frac{\sqrt{R} h(DR)}{R^{\frac{3}{2}(1+\varrho)}} \right) \quad \left(\begin{matrix} DR \equiv 1 \pmod{4} \\ R \text{ rel. pr. zu } 2D \end{matrix} \right)$$

bestimmen. Es möge r eine in der Größe D nicht aufgehende ungerade Primzahl bedeuten. Wir zerlegen dann $\{D\}$, indem wir die Zahlen R nach der höchsten in ihnen enthaltenen Potenz von r ordnen, in eine Summe von Grenzwerten

$$\{D\} = \sum_{s=0}^{\infty} \{D; s\},$$

$$\{D; s\} = \lim \left(\varrho \sum_{R_0} \frac{\sqrt{R_0 r^s} h(D R_0 r^s)}{(R_0 r^s)^{\frac{3}{2}(1+\varrho)}} \right) = \frac{1}{r^s} \lim \left(\varrho \sum_{R_0} \frac{\sqrt{R_0} h(D R_0 r^s)}{R_0^{\frac{3}{2}(1+\varrho)}} \right).$$

$$[Dr^s R_0 \equiv 1 \pmod{4}; R_0 \text{ rel. pr. zu } Dr]$$

Die hier auftretenden einzelnen Grenzwerte

$$\lim \left(\varrho \sum_{R_0} \frac{\sqrt{R_0} h(D r^s R_0)}{R_0^{\frac{3}{2}(1+\varrho)}} \right)$$

sind, wenn $s > 0$ ist, gleich $\{Dr^s\} = \{Dr\}$. Ist aber $s = 0$, so können wir die Summe

$$\{D; 0\} = \lim \left(\varrho \sum_{R_0} \frac{\sqrt{R_0} h(D R_0)}{R_0^{\frac{3}{2}(1+\varrho)}} \right) \quad \left(\begin{matrix} DR_0 \equiv 1 \pmod{4} \\ R_0 \text{ rel. pr. zu } Dr \end{matrix} \right)$$

in zwei Partialsummen L_+ und L_- zerlegen, indem wir alle diejenigen Glieder von $\{D; 0\}$, für welche R_0 quadratischer Rest von r ist, in eine

Summe L_+ und alle diejenigen Glieder von $\{D; 0\}$, für welche R_0 quadratischer Nichtrest von r ist, in eine Summe L_- zusammenfassen. Jeder der beiden Grenzwerte $L_\varepsilon (\varepsilon = +1, -1)$ läßt sich schreiben

$$\begin{aligned} L_\varepsilon &= \frac{1}{1 - \left(\frac{-D}{r}\right) \frac{\varepsilon}{r}} \lim \left(\varrho \sum_{R_0} \frac{\sqrt{R_0} h(D R_0) \left[1 - \left(\frac{-D R_0}{r}\right) \frac{1}{r}\right]}{R_0^{\frac{3}{2}(1+\varrho)}} \right) \\ &= \frac{1}{1 - \left(\frac{-D}{r}\right) \frac{\varepsilon}{r}} \lim \left(\varrho \sum \frac{\sqrt{R_0} h(D r^2 R_0)}{R_0^{\frac{3}{2}(1+\varrho)}} \right); \quad \left[\left(\frac{R_0}{r}\right) = \varepsilon\right] \\ L_\varepsilon &= \frac{1}{2} \cdot \frac{\{D r^2\}}{1 - \left(\frac{-D}{r}\right) \frac{\varepsilon}{r}}, \end{aligned}$$

so daß sich

$$\{D; 0\} = L_+ + L_- = \frac{\{D r\}}{1 - \frac{1}{r^2}}$$

ergibt. Demnach finden wir

$$\{D\} = \{D r\} \left(\frac{1}{1 - \frac{1}{r^2}} + \frac{1}{r} + \frac{1}{r^2} + \dots \right) = \{D r\} \frac{1 - \frac{1}{r^3}}{\left(1 - \frac{1}{r}\right) \left(1 - \frac{1}{r^2}\right)},$$

d. i.

$$\text{für } s \geq 1: \{D r^s\} = \{D r\} = \{D\} \cdot \frac{(r)_1 \cdot (r)_2}{(r)_s}.$$

Durch eine wiederholte Anwendung dieser Formel gewinnen wir die Gleichung

$$(66) \quad \{D\} = \{1\} \cdot \frac{(D)_1 \cdot (D)_2}{(D)_3}.$$

Die Relation

$$\frac{(200')_1 \cdot (200')_2 e_3 M_0}{(200')_3 \cdot 2 S_3} = \{00'\}$$

nimmt jetzt die Form an

$$M_0 = \frac{14}{3 e_3} S_3 \cdot \{1\}.$$

Die Größe M_0 ist mithin eine Konstante; wir bestimmen dieselbe aus dem Maß des Genus $\begin{pmatrix} 1, & 1 \\ 1, & 1 \end{pmatrix}$. Bekanntlich ist dieses Maß gleich $\frac{1}{24}$, und es wird folglich $M_0 = \frac{1}{12}$.

Dieser Wert von M_0 kann auch direkt mit Hilfe der Formel (66) hergeleitet werden. In der Tat, nehmen wir an, daß die Größe D die sämtlichen ungeraden Primzahlen enthält, die unterhalb einer Größe Ω liegen, so werden die Grenzwerte der Ausdrücke $\frac{\{D\}}{(D)_1}, (D)_2, (D)_3$ für $\Omega = \infty$ bzw. gleich $\frac{2}{3\pi} \cdot \frac{1}{4}, \frac{1}{(2)_2 \cdot S_2}, \frac{1}{(2)_3 \cdot S_3}$. Wir finden also $\{1\} = \frac{S_2}{7\pi \cdot S_3}$ und folglich $M_0 = \frac{2 S_2}{3\pi \cdot e_3} = \frac{1}{12}$.

Wir gelangen auf diese Weise zu dem folgenden Resultate:

Das Maß eines Genus

$$\left(\begin{matrix} 1, & 1 \\ o, & o' \end{matrix} \right), C(\varphi), C'(\varphi') \quad [o, o' \equiv 1 \pmod{2}]$$

ist gleich

$$M = \frac{o o'}{12} \cdot \left[1 + \frac{1}{2} (-1)^{\frac{o+1}{2} \cdot \frac{o'+1}{2}} \left(\frac{\varphi}{o} \right) \left(\frac{\varphi'}{o'} \right) \right] \cdot \frac{1}{2^{\vartheta + \vartheta'}} \cdot \prod_p \left[1 + \left(\frac{-o' \varphi}{p} \right) \frac{1}{p} \right] \cdot \prod_{p'} \left[1 + \left(\frac{-o \varphi'}{p'} \right) \frac{1}{p'} \right] \cdot \prod_{r=r'} \left(1 - \frac{1}{rr'} \right).$$

Dieser Satz ist bereits von Eisenstein in Band 35 von Crelles Journal angegeben worden.

Insbesondere schließen wir hieraus für das Maß eines Genus

$$\varphi': \left(\begin{matrix} 1, & 1 \\ 1, & m \end{matrix} \right), \left(\frac{\varphi}{q} \right),$$

wenn die Zahl m ungerade ist und im ganzen μ Primzahlen q enthält, die Gleichung

$$M = \frac{m}{12} \left[1 + \frac{1}{2} (-1)^{\frac{m+1}{2}} \left(\frac{\varphi}{m} \right) \right] \cdot \frac{1}{2^\mu} \prod_q \left[1 + \left(\frac{-\varphi}{q} \right) \frac{1}{q} \right].$$

Mit Hilfe ähnlicher Betrachtungen können wir auch das Maß eines beliebigen Genus von Formen mit drei Variablen bestimmen.

Wir erwähnen hier nur noch den besonderen Fall, daß das Maß eines Genus

$$\varphi': \left(\begin{matrix} 2, & 1 \\ 1, & 2m \end{matrix} \right), \left(\frac{\varphi}{q} \right), (-1)^{\frac{\varphi-1}{2}} = -1 \quad [m \equiv 1 \pmod{2}]$$

gleich

$$\frac{m}{48} \left[1 + (-1)^{\frac{m+1}{2}} \left(\frac{\varphi}{m} \right) \right] \frac{1}{2^\mu} \prod_q \left[1 + \left(\frac{-\varphi}{q} \right) \frac{1}{q} \right]$$

ist.

III. ($n = 4$). Wir werden die Ordnungen

$$O_I'(d): \left(\begin{matrix} 1, & 1, & 1 \\ 1, & 1, & d \end{matrix} \right)$$

untersuchen, welche eine so wichtige Rolle in der Theorie der Darstellung ganzer Zahlen durch eine Summe von fünf Quadraten spielen. — Es möge 2^ν die höchste in d aufgehende Potenz von 2 sein. Setzen wir $d = 2^\nu \cdot d_0$ [$d_0 \equiv 1 \pmod{2}$] und bezeichnen wir mit $p_1, p_2, \dots, p_\vartheta$ die in d_0 enthaltenen ungeraden Primzahlen, mit Φ' eine Grundform der Ordnung $O_I'(d)$ für den Modul $2d$, so besitzt die Form Φ' , wenn φ_1 den

ersten Koeffizienten ihrer adjungierten Form Φ bedeutet, je nach den Fällen $v < 2$, $v = 2$, $v > 2$, die $\vartheta_0 = \vartheta$, $\vartheta + 1$, $\vartheta + 2$ Charaktere

$$\left(\frac{\varphi_1}{p_1}\right), \left(\frac{\varphi_1}{p_2}\right), \dots, \left(\frac{\varphi_1}{p_g}\right), \quad (v < 2)$$

$$C(\varphi_1) \quad \left(\frac{\varphi_1}{p_1}\right), \left(\frac{\varphi_1}{p_2}\right), \dots, \left(\frac{\varphi_1}{p_g}\right), (-1)^{\frac{\varphi_1-1}{2}}, \quad (v = 2)$$

$$\left(\frac{\varphi_1}{p_1}\right), \left(\frac{\varphi_1}{p_2}\right), \dots, \left(\frac{\varphi_1}{p_g}\right), (-1)^{\frac{\varphi_1-1}{2}}, \left(\frac{2}{\varphi_1}\right). \quad (v > 2)$$

Folglich besteht die Ordnung $O'_I(d)$ aus $g = 2^\vartheta$ verschiedenen Genera $G'_1, G'_2, \dots, G'_g$. Wir bezeichnen die Maße dieser Genera durch $M'_1, M'_2, \dots, M'_g$.

Ist p eine der ϑ Primzahlen $p_1, p_2, \dots, p_g$, so besitzt die Kongruenz $\Phi(\xi_i) \equiv 0 \pmod{p}$ p^3 Lösungen $(\xi_i) \pmod{p}$ und die Kongruenz $\Phi(\xi_i) \equiv 0 \pmod{2}$ im ganzen 2^3 Lösungen $(\xi_i) \pmod{2}$. Infolgedessen wird die Anzahl der inkongruenten Wertsysteme $(\xi_i) \pmod{p}$ oder $\pmod{2}$, für welche die Form $\Phi(\xi_i)$ zu einer der Primzahlen p oder zu der Zahl 2 relativ prim wird, gleich $p^4 - p^3 = p^4 \left(1 - \frac{1}{p}\right)$ oder gleich $2^4 - 2^3 = 2^4 \left(1 - \frac{1}{2}\right)$, und die Anzahl der nach dem Modul $2d$ inkongruenten Wertsysteme (ξ_i) , für welche $\Phi(\xi_i)$ einen zu $2d$ relativ primen Wert annimmt, ist gleich $(2d)^4 \cdot (2d)_1$.

Bedeutet jetzt $P(\varphi_1)$ irgendein Produkt der Charaktere $C(\varphi_1)$, so muß die Funktion $P(m)$ (m relativ prim zu $2d$) den Bedingungen

$$P(m) \cdot P(m_0) = P(m \cdot m_0), \quad P(m) \cdot P(m) = 1$$

genügen. Bestimmen wir jetzt ein vollständiges Formensystem $\Phi_1, \Phi_2, \dots, \Phi_K$ für die der Ordnung O'_I adjungierte Ordnung O_I und bilden wir die Summe

$$\sum_{\varrho} = \sum_{k=1}^K \left(\sum_{\{\Phi_k(\xi_1, \xi_2, \xi_3, \xi_4)\}^{\frac{4}{2}(1+\varrho)}} \frac{P_k(\varphi_1) \cdot \{t(\Phi_k)\}^{-1}}{\{ \Phi_k(\xi_1, \xi_2, \xi_3, \xi_4) \}^{\frac{4}{2}(1+\varrho)}} \right)$$

über alle möglichen ganzzahligen Wertsysteme $\xi_1, \xi_2, \xi_3, \xi_4$ ohne gemeinsamen Teiler, für welche der Ausdruck $\Phi_k(\xi_i)$ zu $2d$ relativ prim ausfällt, so wird dieselbe für jedes positive ϱ konvergieren, und der Grenzwert von $\varrho \cdot \sum_{\varrho}$ wird für unendlich abnehmendes ϱ gleich

$$L = \left(\sum_{i=1}^g P^{(\vartheta)}(\varphi_1) M'_i \right) \cdot \frac{(2d)_1 \cdot e_4}{(2d)_4 \cdot S_4 \sqrt{d^3}}$$

werden.

Wir ordnen jetzt die Summe $\sum_{\varrho}$ nach den numerischen Werten der Zahlen $\Phi_k(\xi_1, \xi_2, \xi_3, \xi_4)$. Das Maß der Darstellungen einer zu $2d$ relativ primen Zahl m , in welcher μ ungerade Primzahlen $q_1, q_2, \dots, q_{\mu}$ auf-

gehen, durch die Formen Φ_k ist gleich dem 2^k -fachen Maße des Genus

$$\chi': \begin{pmatrix} 1, & 1 \\ 1, & m \end{pmatrix}, \quad \left(\frac{\chi}{q}\right) = \left(\frac{-d}{q}\right),$$

d. i. gleich

$$\frac{m}{12} \left[1 - \frac{1}{2} (-1)^{\frac{m-1}{2} \cdot \frac{d_0-1}{2}} \left(\frac{m}{d_0}\right) \left(\frac{2^v}{m}\right) \right] \cdot \prod_q \left[1 + \left(\frac{d}{q}\right) \frac{1}{q} \right].$$

Infolgedessen können wir für die Summe $\sum_q$ schreiben

$$\begin{aligned} \sum_q &= \frac{1}{12} \sum_m \left\{ \left[1 - \frac{1}{2} (-1)^{\frac{m-1}{2} \cdot \frac{d_0-1}{2}} \left(\frac{m}{d_0}\right) \left(\frac{2^v}{m}\right) \right] \cdot \prod_q \left[1 + \left(\frac{d}{q}\right) \frac{1}{q} \right] \cdot \frac{m P(m)}{m^{2(1+q)}} \right\} \\ &= \frac{1}{12} T_q - \frac{1}{24} T_q^0, \end{aligned}$$

wo

$$T_q = \sum_m \prod_q \left[1 + \left(\frac{d}{q}\right) \frac{1}{q} \right] \frac{m P(m)}{m^{2(1+q)}}, \quad T_q^0 = \sum_m \prod_q \left[1 + \left(\frac{d}{q}\right) \frac{1}{q} \right] \left(\frac{d}{m}\right) \frac{m P(m)}{m^{2(1+q)}}$$

ist.

Die Summen T_q und T_q^0 können wir mit Hilfe der in I. gegebenen Formel (64) umformen. Da die Größe $P(m)$ dem absoluten Werte nach gleich 1 ist, ergibt sich

$$T_q = \frac{\sum \frac{m P(m)}{m^{2(1+q)}} \cdot \sum \left(\frac{d}{m}\right) \frac{P(m)}{m^{2(1+q)}}}{\sum \frac{1}{m^{4(1+q)}}}, \quad T_q^0 = \frac{\sum \left(\frac{d}{m}\right) \frac{m P(m)}{m^{2(1+q)}} \cdot \sum \frac{P(m)}{m^{2(1+q)}}}{\sum \frac{1}{m^{4(1+q)}}}.$$

Es wird also

$$L = \lim_{q=0} \left(q \sum_q \right) = \frac{1}{12} \lim_{q=0} (q T_q) - \frac{1}{24} \lim_{q=0} (q T_q^0)$$

und

$$\lim_{q=0} (q T_q) = \frac{1}{(2d)_4 \cdot S_4} \cdot \sum \left(\frac{d}{m}\right) \frac{P(m)}{m^2} \cdot \frac{1}{2} \lim_{q=0} \left(q \sum \frac{P(m)}{m^{1+q}} \right),$$

$$\lim_{q=0} (q T_q^0) = \frac{1}{(2d)_4 \cdot S_4} \cdot \sum \frac{P(m)}{m^2} \cdot \frac{1}{2} \lim_{q=0} \left(q \sum \frac{P(m)}{m^{1+q}} \left(\frac{d}{m}\right) \right).$$

Für die Größe $P(m)$ wählen wir der Reihe nach die 2^q Glieder des über alle $\mathfrak{d}_0$ Größen $C(m)$ ausgedehnten Produktes

$$\Pi = \prod (1 + C(m)) = 1 + \dots$$

Durch jedesmalige Vergleichung der beiden Ausdrücke des Grenzwertes L erhalten wir dann 2^q lineare Relationen zwischen den 2^q Größen M'_i , vermittels deren diese Größen eindeutig bestimmt werden können. Denn die Determinante dieser 2^q Gleichungen ist, wie man leicht erkennt, dem absoluten Werte nach gleich $2^{q_0 \cdot 2^{q_0-1}}$.

Die beiden Grenzwerte

$$l = \lim_{\varrho=0} \left(\varrho \sum \frac{P(m)}{m^{1+\varrho}} \right), \quad l_0 = \lim_{\varrho=0} \left(\varrho \sum \frac{P(m) \left(\frac{d}{m} \right)}{m^{1+\varrho}} \right)$$

sind bereits von Dirichlet angegeben worden. In den Fällen $v=0$, $d_0 \equiv -1 \pmod{4}$ und $v=1$ ist der Grenzwert l_0 gleich Null, während der Grenzwert l gleich $(2d)_1$ oder gleich 0 ist, je nachdem $P(m)$ mit dem Gliede 1 des Produktes $\prod$ übereinstimmt oder nicht übereinstimmt. In diesen Fällen $v=0$, $d_0 \equiv -1 \pmod{4}$ und $v=1$ erhalten wir daher

$$M'_i = \frac{1}{24e_4} \sqrt{d^3} \cdot \frac{1}{2^{\varphi_0}} \sum \left(\frac{d}{m} \right) \frac{1}{m^2}.$$

Wenn dagegen $v=0$, $d_0 \equiv 1 \pmod{4}$ oder $v \geq 2$ wird, so erscheint die Größe

$$\left(\frac{d}{m} \right) = \left(\frac{m}{d_0} \right) \cdot \left(\frac{2^v}{m} \right) (-1)^{\frac{m-1}{2} \cdot \frac{d_0-1}{2}}$$

selbst unter den Gliedern des Produktes $\prod$; es findet sich infolgedessen der Grenzwert l_0 gleich $(2d)_1$ oder gleich 0, je nachdem $P(m) = \left(\frac{d}{m} \right)$ ist oder nicht; ähnlich wird der Grenzwert l gleich $(2d)_1$ oder gleich 0, je nachdem $P(m)$ mit 1 übereinstimmt oder nicht übereinstimmt. Wir bekommen demnach für die Fälle $v=0$, $d_0 \equiv 1 \pmod{4}$ oder $v \geq 2$ die Formeln

$$M'_i = \frac{1}{24e_4} \left[1 - \frac{1}{2} (-1)^{\frac{\varphi_1-1}{2} \cdot \frac{d_0-1}{2}} \left(\frac{\varphi_1}{d_0} \right) \left(\frac{2^v}{\varphi_1} \right) \right] \sqrt{d^3} \cdot \frac{1}{2^{\varphi_0}} \sum \left(\frac{d}{m} \right) \frac{1}{m^2}.$$

Mit Hilfe dieser Gleichungen für die Größen M'_i können wir insbesondere das Maß des in Kap. XXI betrachteten Genus $G'_i(d)$ finden. Wir bekommen so den folgenden Satz:

Die Anzahl der eigentlichen Darstellungen einer ganzen Zahl d durch eine Summe von fünf Quadraten, welche nicht sämtlich ungerade sind, ist, wenn $d \equiv 3 \pmod{4}$ oder $\equiv 2 \pmod{4}$ ist, gleich

$$1920 \cdot \frac{1}{12\pi^2} \sqrt{d^3} \cdot \sum \left(\frac{d}{m} \right) \frac{1}{m^2},$$

dagegen, wenn $d \equiv 1 \pmod{4}$ oder $\equiv 0 \pmod{4}$ ist, gleich

$$\frac{1}{2} \cdot 1920 \cdot \frac{1}{12\pi^2} \sqrt{d^3} \cdot \sum \left(\frac{d}{m} \right) \frac{1}{m^2}.$$

Wir können diese beiden Formeln in eine einzige zusammenfassen und erhalten dann den folgenden Satz:

Die Anzahl der eigentlichen Darstellungen einer ganzen Zahl d durch eine Summe von fünf nicht lauter ungeraden Quadraten beträgt

$$\frac{40}{\pi^2} \left(3 - (-1)^{\left[\frac{d}{2} \right]} \right) \sqrt{d^3} \cdot \sum \left(\frac{d}{m} \right) \frac{1}{m^2}. \quad (m \text{ rel. pr. zu } 2d)$$

Ähnlich können wir die Maße der Genera einer Ordnung

$$O'_{II}(d): \begin{pmatrix} 2, & 1, & 2 \\ 1, & 1, & d \end{pmatrix}$$

bestimmen. Die Formeln für das Maß des in Kap. XXI betrachteten Genus $G'_{II}(d)$ ergeben dann den folgenden Satz:

Eine Zahl $d \equiv 5 \pmod{8}$ besitzt

$$\frac{32}{\pi^2} \sqrt{d^3} \cdot \sum \left(\frac{d}{m} \right) \frac{1}{m^2} \quad (m \text{ rel. pr. zu } 2d)$$

eigentliche Darstellungen durch eine Summe von fünf ungeraden Quadraten. Eine Zahl d , welche nicht kongruent $5 \pmod{8}$ ist, läßt sich nicht durch eine Summe von fünf ungeraden Quadraten darstellen.

Die in diesen Formeln auftretenden Summen $\frac{1}{\pi^2} \sqrt{d^3} \cdot \sum \left(\frac{d}{m} \right) \frac{1}{m^2}$ sind spezielle Fälle der Summen

$$\sum_s = \frac{\sqrt{d^{2s-1}}}{\pi^s} \sum_m \left(\frac{(-1)^s \cdot d}{m} \right) \frac{1}{m^s} \quad (d > 0; m \text{ rel. pr. zu } 2d)$$

Indem wir, je nachdem $s = 2s_0$ oder $s = 2s_0 - 1$ ist, von den bekannten Werten der Reihe $\sum_{k=1}^{\infty} \frac{\cos 2\pi k z}{k^{2s_0}}$ oder der Reihe $\sum_{k=1}^{\infty} \frac{\sin 2\pi k z}{k^{2s_0-1}}$ Gebrauch machen, können wir für diese Summen analoge Ausdrücke finden, wie sie Dirichlet für den Fall $s = 1$ aufgestellt hat.

Kap. XXIII. Maß eines beliebigen Genus einer Ordnung

$$\begin{pmatrix} 1 \\ o_h \end{pmatrix} [o_h \equiv 1 \pmod{2}].$$

Die Untersuchung über die Maße positiver Genera habe ich soweit gefördert, daß ich in kurzer Zeit das Maß eines jeden beliebigen Genus angeben zu können hoffe. Um einen Begriff von den Resultaten zu geben, welche ich auf diesem interessanten Gebiet erhalten habe, will ich das Maß eines Genus G :

$$(o_0 = 0) \quad \begin{pmatrix} 1, & 1, & \dots, & 1, & 1 \\ o_1, & o_2, & \dots, & o_{n-2}, & o_{n-1} \end{pmatrix} \quad (o_n = 0)$$

$$C(\varphi_1), C(\varphi_2), \dots, C(\varphi_{n-2}), C(\varphi_{n-1})$$

mitteilen, für welches die Invarianten o_h sämtlich ungerade sind.

Sind $u', u'', \dots; v', v'', \dots$ irgend zwei Reihen ganzer Zahlen, so wollen wir durch das Symbol $NP(\bar{u}', \bar{u}'', \dots; \bar{v}', \bar{v}'', \dots)$ alle Primzahlen bezeichnen, welche in den sämtlichen Primzahlen $u', u'', \dots$ enthalten sind

und welche gleichzeitig zu den sämtlichen Zahlen $v', v'', \dots$ relativ prim ausfallen.

Es sei τ_h die Anzahl der verschiedenen Zahlen $NP(o_h^+)$. Die Anzahl aller der Ordnung $\binom{1}{o_h}$ angehörigen verschiedenen Genera wird dann

$$\sum_{h=1}^{n-1} \tau_h$$

gleich $g = 2^{h=1}$ sein.

Wir wollen mit k die sämtlichen Zahlen der Reihe $1, 2, \dots, \left[\frac{n-1}{2}\right]$ bezeichnen und mit i alle Zahlen, für welche $k \leq i \leq n-k$ ist. Von den Zahlen o_{i-k} und o_{i+k} ist dann mindestens eine von Null verschieden, und es sind die Größen

$$r_i^{(k)} = o_i^k \cdot \prod_{h=1}^{k-1} (o_{i-h} \cdot o_{i+h})^{k-h}$$

gleichfalls von Null verschieden. — Wir wollen die Zahlen $NP(o_{i-k}^+, o_{i+k}^+; r_i^{(k)})$ durch $\omega_i^{(k)}$ bezeichnen, ferner, wenn $k < i < n-k$ ist, die Zahlen $NP(o_{i-k}^+, o_{i+k}^+; r_i^{(k)})$, wenn aber $i-k=0$ oder $i+k=n$ ist, die Zahlen $NP(o_{i-k}^+, o_{i+k}^+; r_i^{(k)})$ durch $\vartheta_i^{(k)}$ bezeichnen. Die Zahlen $NP\left(\prod_{h=1}^{2k} o_h^+\right)$ mögen $p_0^{(k)}$ und die Zahlen $NP\left(\prod_{h=1}^{2k} o_{n-h}^+\right)$ mögen $p_n^{(k)}$ heißen. Bilden wir das Produkt

$$\prod \frac{\left(1 - \frac{1}{(p_0^{(k)})^{2k}}\right) \left(1 - \frac{1}{(p_n^{(k)})^{2k}}\right)}{\left(1 - \frac{1}{(\omega_i^{(k)})^{2k}}\right) \left(1 - \frac{1}{(\vartheta_i^{(k)})^{2k}}\right)}$$

über alle möglichen Primzahlen

$$\omega_i^{(k)}, \vartheta_i^{(k)}, p_0^{(k)}, p_n^{(k)}, \quad (k \leq i \leq n-k)$$

welche einem bestimmten k entsprechen, so wird dasselbe, wie man leicht erkennt, ein vollständiges Quadrat, und wir können es durch Π_k^2 bezeichnen ($\Pi_k > 0$). Das Produkt

$$\prod \left[1 + \left(\frac{(-1)^k \cdot r_i^{(k)} \cdot \varphi_{i-k} \cdot \varphi_{i+k}}{\omega_i^{(k)}} \right) \frac{1}{(\omega_i^{(k)})^k} \right],$$

welches über alle Primzahlen $\omega_i^{(k)}$ ausgedehnt sei, die einem bestimmten k entsprechen, möge gleich D_k gesetzt werden. Ferner schreiben wir

$$\prod_{k=1}^{\left[\frac{n-1}{2}\right]} \Pi_k = \Pi, \quad \prod_{k=1}^{\left[\frac{n-1}{2}\right]} D_k = D$$

und

$$\prod_{h=1}^{n-1} o_h^{h(n-h)} = \Delta.$$

Endlich führen wir die beiden Einheiten ein

$$\delta_0 = \prod_{h=1}^{n-1} \left(\frac{\varphi_h}{o_h} \right) \cdot (-1)^{\left[\frac{n}{4} \right] + \sum_{t' < t''}^{1, n-1} \frac{\left(\prod_{h=1}^{t'} o_h \right)^{-1}}{2} \cdot \frac{\left(\prod_{h=1}^{t''} o_h \right)^{-1}}{2}} \cdot (-1)^{\left[\frac{n}{2} \right] \left(\left[\frac{n}{2} \right] + \sum_{t=1}^{n-1} \frac{\left(\prod_{h=1}^t o_h \right)^{-1}}{2} \right)},$$

$$\delta_n = \prod_{h=1}^{n-1} \left(\frac{\varphi_{n-h}}{o_{n-h}} \right) \cdot (-1)^{\left[\frac{n}{4} \right] + \sum_{t' < t''}^{1, n-1} \frac{\left(\prod_{h=1}^{t'} o_{n-h} \right)^{-1}}{2} \cdot \frac{\left(\prod_{h=1}^{t''} o_{n-h} \right)^{-1}}{2}} \cdot (-1)^{\left[\frac{n}{2} \right] \left(\left[\frac{n}{2} \right] + \sum_{t=1}^{n-1} \frac{\left(\prod_{h=1}^t o_{n-h} \right)^{-1}}{2} \right)}.$$

Wenn $n \equiv 1 \pmod{2}$ oder $n \equiv 0 \pmod{2}$, $\Delta \equiv (-1)^{\frac{n}{2}} \pmod{4}$ ist, findet man $\delta_0 = \delta_n$, und wir setzen

$$\delta = \frac{\delta_0 + \delta_n}{2},$$

dagegen, wenn $n \equiv 0 \pmod{2}$, $\Delta \equiv -(-1)^{\frac{n}{2}} \pmod{4}$ ist, $\delta = 0$.

Alsdann gelten für das Maß M des Genus G die beiden Formeln:

$$\text{für } n \equiv 1 \pmod{2}: \quad M = \frac{\varepsilon_n}{g} \cdot II \cdot D \cdot \Delta^{\frac{1}{2}} \left(1 + \frac{\delta}{2^{\frac{n-1}{2}}} \right)$$

und für $n \equiv 0 \pmod{2}$:

$$M = \frac{\varepsilon_n}{g} \cdot II \cdot D \cdot \Delta^{\frac{1}{2}} \left(1 + \frac{\delta}{2^{\frac{n}{2}-1}} \right) \cdot \pi^{-\frac{n}{2}} \sum \left(\frac{(-1)^{\frac{n}{2}} \Delta}{m} \right) \frac{1}{m^{\frac{n}{2}}},$$

worin die Größen ε_n gewisse rationale Konstante bedeuten, welche nur von der Zahl n abhängen. —

Um diese Formeln für den Fall n zu bestätigen, falls sie für den Fall $n-1$ bereits bewiesen sind, können wir, wenn $n \equiv 0 \pmod{2}$ ist, genau auf dieselbe Art verfahren, wie Dirichlet zur Bestimmung des Maßes eines Genus mit zwei Variablen getan hat, und, wenn $n \equiv 1 \pmod{2}$ ist, genau auf dieselbe Art, wie wir soeben zur Aufstellung des Maßes eines Genus mit drei Variablen verfahren sind.

Note über die Kongruenzen $f \equiv g \pmod{q'}$.

In dieser Note wollen wir einen Beweis des Satzes M., Kap. X, geben. Wir bezeichnen durch $q^{\partial k-1(f)}$ die höchste Potenz einer Primzahl q ,

welche in den sämtlichen k -reihigen Unterdeterminanten einer Form f aufgeht, und wir setzen $q^{\partial_k - \partial_{k-1}} = q^{\nu_k}$, $q^{\nu_k - \nu_{k-1}} = q^{\omega_k}$.

I. [$q = p \equiv 1 \pmod{2}$.] — Zwei Klassen f und g von Formen mit n Variablen sind in bezug auf einen Modul $p^t (> p^{\nu_{n-1}(f)}, > p^{\nu_{n-1}(g)})$ kongruent, wenn die Relationen

$$(67) \quad f\{m; p^t\} = g\{m; p^t\} \quad [m \equiv 1, 2, \dots, p^t \pmod{p^t}]$$

statthaben und wenn die Determinanten $\Delta(f)$ und $\Delta(g)$ nach dem kleineren der beiden Moduln $p^{t+\partial_{n-2}(f)}$, $p^{t+\partial_{n-2}(g)}$ kongruent sind.

Beweis. Wir können voraussetzen, daß die eine der beiden Formen f und g in bezug auf p primitiv ist. Denn ist $f = p^\partial \cdot f_0$, $g = p^\partial \cdot g_0$ ($\partial > 0$, $\partial \leq t$), so gilt $f(h; p^t) = f_0(h \cdot p^\partial; p^t) = p^{n\partial} \cdot f_0(h; p^{t-\partial})$ und $g(h; p^t) = p^{n\partial} \cdot g_0(h; p^{t-\partial})$; die Beziehungen (67) ergeben also

$$f_0(h; p^{t-\partial}) = g_0(h; p^{t-\partial});$$

andererseits liefert die Kongruenz $f_0 \equiv g_0 \pmod{p^{t-\partial}}$ sofort $f \equiv g \pmod{p^t}$.

Ist die Form f primitiv in bezug auf p , so wird die Kongruenz $f(x_i) \equiv \alpha \pmod{p^t}$ für gewisse zu p relativ prime Zahlen α lösbar sein. Für diese selben Zahlen α hat dann die Kongruenz $g(y_i) \equiv \alpha \pmod{p^t}$ gleichfalls Lösungen; denn wir haben $g\{\alpha; p^t\} = f\{\alpha; p^t\}$ vorausgesetzt. Folglich muß die Form g ebenso wie f in bezug auf p primitiv sein. — Da die Zahlen α zu p relativ prim sind, können die x_i in einer Kongruenz $f(x_i) \equiv \alpha \pmod{p^t}$ nicht sämtlich durch p teilbar sein. Es ist demnach möglich, n Zahlen $\xi_i \equiv x_i \pmod{p^t}$ zu bestimmen, deren größter gemeinsamer Teiler gleich 1 ist. Alsdann kann man eine Substitution

$$S = \begin{pmatrix} \xi_1, & \dots, & \dots, & \dots \\ \xi_2, & \dots, & \dots, & \dots \\ \dots & \dots & \dots & \dots \\ \xi_n, & \dots, & \dots, & \dots \end{pmatrix}$$

von der Determinante 1 finden, welche die Form f in eine Form $\widehat{f}_{(1)}$ überführt, deren erster Koeffizient kongruent $\alpha \pmod{p^t}$ wird. Nach Kap. II läßt sich diese Form $\widehat{f}_{(1)}$ noch in einen Repräsentanten von der Gestalt

$$f_{(1)} \equiv \alpha \xi^2 + F^{(1)} \pmod{p^t}$$

transformieren. Für die Form $f_{(1)}$ gelten die Beziehungen

$$\partial_{k-1}(F^{(1)}) = \partial_k(f), \quad (k \geq 1); \quad \Delta(F^{(1)}) \equiv \frac{\Delta(f)}{\alpha} \pmod{p^{t+\partial_{n-2}(f)}}$$

und

$$F^{(1)}(h; p^t) = \frac{f(h; p^t)}{(\alpha h; p^t)}, \quad (\alpha h; p^t) \neq 0.$$

* Siehe Kapitel VII.

Die nämlichen Schlüsse finden auf die Form g Anwendung. Wir erhalten für diese Form einen Repräsentanten

$$g_{(1)} \equiv \alpha \xi^2 + G^{(1)} \pmod{p^t},$$

für welchen

$$\partial_{k-1}(G^{(1)}) = \partial_k(g), \quad (k \geq 1); \quad \Delta(G^{(1)}) \equiv \frac{\Delta(g)}{\alpha} \pmod{p^{t+\partial_{n-2}(g)}}$$

und

$$G^{(1)}(h; p^t) = \frac{g(h; p^t)}{(\alpha h; p^t)}$$

ist.

Setzen wir jetzt voraus, unser Satz sei bereits für Formen mit $n-1$ Variablen als richtig erkannt, so wird

$$F^{(1)} \simeq G^{(1)} \pmod{p^t}$$

und folglich $f_{(1)} \simeq g_{(1)} \pmod{p^t}$. Damit wird unser Satz also auch für Formen mit n Variablen bewiesen sein.

II. ($g = 2$.) *Zwei Klassen f und g von Formen mit n Variablen sind in bezug auf einen Modul $2^t (> 2^{v_{n-1}(f)}, > 2^{v_{n-1}(g)})$ kongruent, wenn die Relationen*

$$f\{m; 2^t\} = g\{m; 2^t\} \quad [m \equiv 1, 2, \dots, 2^t \pmod{2^t}]$$

und die Kongruenzen

$$(68) \quad 2^{\omega_k(f)} \equiv 2^{\omega_k(g)} \pmod{2} \quad (k = 0, 1, \dots, n-1)$$

statthaben und wenn die Determinanten $\Delta(f)$ und $\Delta(g)$ nach dem kleineren der beiden Moduln $\sigma_{n-1}(f) \cdot 2^{t+\partial_{n-2}(f)}$, $\sigma_{n-1}(g) \cdot 2^{t+\partial_{n-2}(g)}$ kongruent sind.

Beweis. Es genügt, den Fall zu betrachten, daß die eine, f , der beiden Formen primitiv in bezug auf 2 ist.

1. Es sei zunächst $\sigma_1(f) = 1$. Dann können wir ungerade Zahlen α finden, für welche $f\{\alpha; 2^t\} > 0$ ist, und für dieselben Zahlen α wird dann $g\{\alpha; 2^t\} > 0$ sein. Infolgedessen ist die Form g primitiv in bezug auf 2, und es gilt $\sigma_1(g) = 1$.

Sobald für eine ungerade Zahl α die Kongruenzen

$$f(\xi_i) \equiv \alpha \pmod{2^t}, \quad g(\eta_i) \equiv \alpha \pmod{2^t}$$

lösbar sind, können wir f und g in zwei Repräsentanten

$$f_{(1)} \equiv \alpha \xi^2 + F^{(1)} \equiv \alpha \xi^2 + 2^{\omega_1(f)} \cdot f^{(1)} \pmod{2^t},$$

$$g_{(1)} \equiv \alpha \xi^2 + G^{(1)} \equiv \alpha \xi^2 + 2^{\omega_1(g)} \cdot g^{(1)} \pmod{2^t}$$

transformieren, in denen $f^{(1)}$ und $g^{(1)}$ primitive Formen in bezug auf 2 vorstellen.

Wenn eine der beiden Zahlen $\omega_1(f)$, $\omega_1(g)$ gleich Null ist, so verschwindet die andere gleichfalls; denn es ist $2^{\omega_1(f)} \equiv 2^{\omega_1(g)} \pmod{2}$. In dem Falle, daß $2^{\omega_1(f)} = 2^{\omega_1(g)} = 1$ ist, kann man stets voraussetzen, daß die ungerade Zahl α und die Systeme $\xi_i \pmod{2^t}$ und $\eta_i \pmod{2^t}$ so gewählt

sind, daß die Formen $f^{(1)}$ und $g^{(1)}$ alle beide eine erste Invariante σ gleich 1 besitzen.

In der Tat, die Form f möge von der in Kap. III aufgestellten Gestalt $f_{(x_1)}$ sein, und es möge

$$S \equiv \begin{pmatrix} \xi_1, & \vartheta_1^1, & \dots, & \vartheta_1^{n-1} \\ \xi_2, & \vartheta_2^1, & \dots, & \vartheta_2^{n-1} \\ \cdot & \cdot & \cdot & \cdot \\ \xi_n, & \vartheta_n^1, & \dots, & \vartheta_n^{n-1} \end{pmatrix} \pmod{2^t}$$

diejenige Substitution bedeuten, welche f in $f_{(1)}$ verwandelt. Ist $\omega_1(f) = 0$, so wird $\kappa_1(f) > 1$. Damit die erste Invariante σ der Form $f^{(1)}$ gleich 2 wird, müssen die Kongruenzen

$$\begin{aligned} \xi_1 + \xi_2 + \dots + \xi_{x_1} &\equiv 1, & \xi_1 \vartheta_1^i + \xi_2 \vartheta_2^i + \dots + \xi_{x_1} \vartheta_{x_1}^i &\equiv 0, \\ \vartheta_1^i + \vartheta_2^i + \dots + \vartheta_{x_1}^i &\equiv 0 \pmod{2} \end{aligned}$$

gelten. Wie man leicht erkennt, können diese Kongruenzen nur in dem Falle bestehen, daß

$$\kappa_1 \equiv 1 \pmod{2}; \quad \xi_1 \equiv \xi_2 \equiv \dots \equiv \xi_{x_1} \equiv 1 \pmod{2}$$

ist. Fällt also $\kappa_1 \equiv 0 \pmod{2}$ aus, so wird stets $\sigma_1(f^{(1)}) = 1$. Ist aber $\kappa_1 \equiv 1 \pmod{2}$ und $\kappa_1 > 1$, so finden sich unter den $2^{n t - 1}$ Systemen $\xi_i \pmod{2^t}$, welche die Kongruenz $f(\xi_i) \equiv 1 \pmod{2}$ erfüllen, $2^{n t - 1} \left(1 - \frac{1}{2^{x_1 - 1}}\right)$ Systeme $\xi_i \pmod{2^t}$, denen eine Form $f^{(1)}$ mit einer Invariante $\sigma_1(f^{(1)}) = 1$ entspricht, und $2^{n t - 1} \cdot \frac{1}{2^{x_1 - 1}}$ Systeme $\xi_i \pmod{2^t}$, denen eine Form $f^{(1)}$ mit einer Invariante $\sigma_1(f^{(1)}) = 2$ entspricht.

Die Kongruenzen (68) ergeben die Beziehung $\kappa_1(f) = \kappa_1(g) = \kappa_1$, und wir gelangen für die Form g zu Resultaten, welche denen für die Form f erhaltenen ganz analog sind. Im Falle $\kappa_1 \equiv 0 \pmod{2}$ findet sich demnach unsere Annahme verwirklicht. Prüfen wir jetzt den Fall, daß $\kappa_1 \equiv 1 \pmod{2}$ und $\kappa_1 > 1$ ist. Wenn keine der Zahlen α , die man mit Hilfe von Systemen $\xi_i \pmod{2^t}$, für welche $\sigma_1(f^{(1)}) = 1$ wird, erhält, mit Hilfe von Systemen $\eta_i \pmod{2^t}$ erhalten werden könnte, für welche $\sigma_1(g^{(1)}) = 1$ ist, so würde man aus den Gleichungen $f\{\alpha; 2^t\} = g\{\alpha; 2^t\}$ schließen, daß mindestens $2^{n t - 1} \left(1 - \frac{1}{2^{x_1 - 1}}\right)$ Systeme $\eta_i \pmod{2^t}$ Formen $g^{(1)}$ mit einer Invariante $\sigma_1(g^{(1)}) = 2$ liefern müßten. Aber das kann nicht sein, da wegen $\kappa_1 > 1$ und $\equiv 1 \pmod{2}$ stets $2^{n t - 1} \cdot \frac{1}{2^{x_1 - 1}} < 2^{n t - 1} \left(1 - \frac{1}{2^{x_1 - 1}}\right)$ ist. Infolgedessen gibt es in der Tat Zahlen α und Systeme $\xi_i \pmod{2^t}$ und $\eta_i \pmod{2^t}$, für welche man $\sigma_1(f^{(1)}) = 1$ und $\sigma_1(g^{(1)}) = 1$ bekommt.

Nach Kap. II gelten jetzt die folgenden Relationen:

$$\begin{aligned}\omega_{k-1}(F^{(1)}) &= \omega_k(f), & \sigma_{k-1}(F^{(1)}) &= \sigma_k(f), \\ \Delta(F^{(1)}) &\equiv \frac{\Delta(f)}{\alpha} \pmod{\sigma_{n-1}(f) \cdot 2^{t+\partial_{n-2}(f)}}, \\ \omega_{k-1}(G^{(1)}) &= \omega_k(g), & \sigma_{k-1}(G^{(1)}) &= \sigma_k(g), \\ \Delta(G^{(1)}) &\equiv \frac{\Delta(g)}{\alpha} \pmod{\sigma_{n-1}(g) \cdot 2^{t+\partial_{n-2}(g)}},\end{aligned}$$

und wir finden die Formeln

$$F^{(1)}(h; 2^t) = \frac{f(h; 2^t)}{(\alpha h; 2^t)}, \quad G^{(1)}(h; 2^t) = \frac{g(h; 2^t)}{(\alpha h; 2^t)}, \quad (\alpha h; 2^t) \neq 0$$

in denen h irgendeine Zahl bedeutet, die inkongruent $2^{t-1} \pmod{2^t}$ ist, und die Formeln

$$\begin{aligned}F^{(1)}(h; 2^t) &= 0, & G^{(1)}(h; 2^t) &= 0, & (\kappa_1 = 1) \\ F^{(1)}(h; 2^t) &= 2^{n_t}, & G^{(1)}(h; 2^t) &= 2^{n_t}, & (\kappa_1 > 1)\end{aligned}$$

wenn $h \equiv 2^{t-1} \pmod{2^t}$ ist. Nehmen wir also an, unser Satz sei bereits für Formen mit $n-1$ Variablen bewiesen, so erhalten wir $F^{(1)} \simeq G^{(1)} \pmod{2^t}$ und daraus $f_{(1)} \simeq g_{(1)} \pmod{2^t}$.

2. Es möge jetzt $\sigma_1(f) = 2$ sein. Da f in bezug auf 2 primitiv ist, wird die Zahl $2^{\omega_0(f)}$ gleich 1, und die Kongruenz $2^{\omega_0(f)} \equiv 2^{\omega_0(g)} \pmod{2}$ ergibt $2^{\omega_0(g)} = 1$. Also ist die Form g gleichfalls primitiv in bezug auf 2. Wir finden $\sigma_1(g) = 2$; denn wäre $\sigma_1(g) = 1$, so erhielten wir $g(2^{t-1}; 2^t) = 0$, während $f(2^{t-1}; 2^t)$ gleich 2^{n_t} wäre. Die Kongruenzen (68) ergeben die Beziehung $\kappa_1(f) = \kappa_1(g)$, welche uns sofort zeigt, daß $f \simeq g \pmod{2}$ ist. Es sei also $t > 1$.

Sobald für eine Zahl $2\alpha \equiv 2 \pmod{4}$ die Kongruenz $f(\xi_i) \equiv 2\alpha \pmod{2^t}$ lösbar ist, können wir f vermittels einer Substitution

$$S \equiv \begin{pmatrix} \xi_1, & \vartheta_1^1, & \dots, & \vartheta_1^{n-1} \\ \xi_2, & \vartheta_2^1, & \dots, & \vartheta_2^{n-1} \\ \dots & \dots & \dots & \dots \\ \xi_n, & \vartheta_n^1, & \dots, & \vartheta_n^{n-1} \end{pmatrix} \pmod{2^t}$$

von der Determinante 1 in eine Form

$$\hat{f}_{(2)} \equiv \begin{pmatrix} 2\alpha, & E_1, & \dots, & E_{n-1} \\ E_1, & \dots, & \dots, & \dots \\ \dots & \dots & \dots & \dots \\ E_{n-1}, & \dots, & \dots, & \dots \end{pmatrix} \pmod{2^t}$$

transformieren. Wir nennen das System $\xi_i \pmod{2^t}$ primär, wenn unter den Zahlen $E_1, \dots, E_{n-1}$ ungerade vorkommen. Allemal, wenn das System ξ_i primär ausfällt, läßt sich die Form $\hat{f}_{(2)}$ in einen Repräsentanten von der Gestalt

$$f_{(2)} \equiv 2(\alpha \xi^2 + \mathfrak{A} \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2) + 2^{\omega_2(f)} \cdot f^{(2)} \equiv \chi + F^{(2)} \pmod{2^t}$$

verwandeln, in dem $\mathfrak{A} \equiv 1 \pmod{2}$ ist.

Wir wollen jetzt voraussetzen, die Form f sei von der in Kap. III betrachteten Gestalt $f_{(x_1)}$. Damit die Zahlen $E_1, \dots, E_{n-1}$ sämtlich gerade sind, müssen dann die Kongruenzen

$$\xi_1 \vartheta_2^i + \xi_2 \vartheta_1^i + \dots + \xi_{x_1-1} \vartheta_{x_1}^i + \xi_{x_1} \vartheta_{x_1-1}^i \equiv 0 \pmod{2}$$

gelten, deren einzige Lösung

$$[\kappa_1 \equiv 0 \pmod{2}], \quad \xi_1 \equiv \xi_2 \equiv \dots \equiv \xi_{x_1} \equiv 0 \pmod{2}$$

ist. Aber diese Kongruenzen sind nur dann mit der Kongruenz $2\alpha \equiv 2 \pmod{4}$ verträglich, wenn $\sigma_{x_1+1}(f) \cdot 2^{\omega_{x_1}(f)} = 2$ ist. Also werden, falls $\sigma_{x_1+1}(f) \cdot 2^{\omega_{x_1}(f)} > 2$ ist, die sämtlichen Systeme ξ_i , für welche $f(\xi_i) \equiv 2 \pmod{4}$ ausfällt, primär sein. Sobald hingegen $\sigma_{x_1+1}(f) \cdot 2^{\omega_{x_1}(f)} = 2$ ist, finden wir unter den 2^{n-1} Systemen $\xi_i \pmod{2^t}$, für welche $f(\xi_i) \equiv 2 \pmod{4}$ ist, $2^{n-1} \left(1 - \frac{1}{2^{x_1}}\right)$ primäre Systeme und $2^{n-1} \cdot \frac{1}{2^{x_1}}$ nichtprimäre Systeme. Dieselben Schlüsse finden auf die Form g Anwendung.

Indem wir für den Fall $\sigma_{x_1+1}(f) \cdot 2^{\omega_{x_1}(f)} = 2$ bemerken, daß die Zahl $2^{n-1} \cdot \frac{1}{2^{x_1}} < 2^{n-1} \left(1 - \frac{1}{2^{x_1}}\right)$ ist, schließen wir wie in 1., daß wir die Zahl 2α so wählen können, daß einerseits die Kongruenz $f(\xi_i) \equiv 2\alpha \pmod{2^t}$ eine primäre Lösung ξ_i und andererseits die Kongruenz $g(\eta_i) \equiv 2\alpha \pmod{2^t}$ eine primäre Lösung η_i besitzt. Ist dies geschehen, so liefert die Form f einen Repräsentanten von der Form $f_{(2)}$, und die Form g kann in eine Form von der Gestalt

$$g_{(2)} \equiv 2(\alpha \xi^2 + A \xi \tilde{\xi} + a \tilde{\xi}^2) + 2\omega_2(g) \cdot g^{(2)} \pmod{2^t}$$

verwandelt werden, in der $A \equiv 1 \pmod{2}$ ist.

Wir können jetzt voraussetzen, daß $a \equiv \tilde{a} \pmod{2}$ ist. In der Tat, es sei zunächst $\sigma_1(f^{(2)}) \cdot 2^{\omega_2(f)} \geq 4$. Alsdann wird die Anzahl der Lösungen der Kongruenz $f_{(2)} \equiv 2 \pmod{4}$ gleich $2^{2n-1} \left[1 - \frac{1}{2} \cdot \left(\frac{2}{4\alpha\tilde{a} - 2}\right)\right]$. Die Kongruenz $2\omega_2(f) \equiv 2\omega_2(g) \pmod{2}$ ergibt $2\omega_2(g) \geq 2$, und es gilt $\sigma_1(g^{(2)}) \cdot 2^{\omega_2(g)} \geq 4$. Denn wäre $\sigma_1(g^{(2)}) \cdot 2^{\omega_2(g)} = 2$, so bekämen wir $2\omega_2(g) = 2$, $\sigma_1(g^{(2)}) = 1$, und die Kongruenz $g_{(2)} \equiv 2 \pmod{4}$ hätte 2^{2n-1} Lösungen, so daß $g_{(2)}\{2; 4\} \neq f_{(2)}\{2; 4\}$ wäre. Ist jetzt $\sigma_1(g^{(2)}) \cdot 2^{\omega_2(g)} \geq 4$, so liefert die Beziehung $f_{(2)}\{2; 4\} = g_{(2)}\{2; 4\}$ sofort $\left(\frac{2}{4\alpha\tilde{a} - 2}\right) = \left(\frac{2}{4\alpha a - A}\right)$, d. i. $\tilde{a} \equiv a \pmod{2}$.

Gilt aber $\sigma_1(f^{(2)}) \cdot 2^{\omega_2(f)} = \sigma_1(g^{(2)}) \cdot 2^{\omega_2(g)} = 2$ und $a \equiv \tilde{a} + 1 \pmod{2}$, so können wir die Form $2\omega_2(g) \cdot g^{(2)}$ zunächst derart transformieren, daß einer ihrer mittleren Koeffizienten, $r_{ii}^{(2)}$, kongruent $2 \pmod{4}$ wird. Durch Anwendung einer Substitution

$$\tilde{\xi} = \tilde{\xi}', \quad x_i^{(2)} = \tilde{\xi}' + x_i'^{(2)}$$

gelangen wir dann zu einer Form, in welcher der Koeffizient a durch

einen Koeffizienten $\alpha_0 \equiv a + 1 \equiv \tilde{a} \pmod{2}$ ersetzt ist, und diese Form kann wiederum in eine Form von der Gestalt $g_{(2)}$ verwandelt werden, in der der Koeffizient α_0 der nämliche ist.

Demnach ist es stets möglich, anzunehmen, daß

$$a \equiv \tilde{a} \pmod{2},$$

d. i.

$$4\alpha\tilde{a} - \mathfrak{A}^2 \equiv 4\alpha a - A^2 \pmod{8}$$

gilt. Infolge dieses Umstandes können wir eine Zahl Z finden, welche der Kongruenz

$$4\alpha\tilde{a} - \mathfrak{A}^2 \equiv (4\alpha a - A^2)Z^2 \pmod{2^{t+1}}$$

genügt, und $g_{(2)}$ geht vermittels einer Substitution

$$\tilde{\xi} \equiv Z \cdot \xi', \quad x_i^{(2)} \equiv \frac{1}{Z} \cdot x_i'^{(2)} \pmod{2^t}$$

von der Determinante 1 in eine Form über, welche dieselbe Gestalt besitzt, in der aber der Rest

$$\begin{pmatrix} 2\alpha, & A \\ A, & 2a \end{pmatrix} \pmod{2^t}$$

durch den Rest

$$\begin{pmatrix} 2\alpha, & AZ \\ AZ, & 2aZ^2 \end{pmatrix} \pmod{2^t}$$

ersetzt ist. Dieser letztere Rest verwandelt sich noch durch eine Substitution

$$S_0 \equiv \begin{pmatrix} 1, & \frac{\mathfrak{A} - AZ}{2\alpha} \\ 0, & 1 \end{pmatrix} \pmod{2^t}$$

in einen Rest

$$\equiv \begin{pmatrix} 2\alpha, & \mathfrak{A} \\ \mathfrak{A}, & 2\tilde{a} \end{pmatrix} \equiv \chi \pmod{2^t}.$$

Für die Form $g_{(2)}$ erhalten wir demnach eine Kongruenz von der Gestalt

$$g_{(2)} \simeq \chi + G^{(2)} \pmod{2^t}.$$

Wir haben jetzt die folgenden Beziehungen:

$$\begin{aligned} \omega_{k-2}(F^{(2)}) &= \omega_k(f), & \sigma_{k-2}(F^{(2)}) &= \sigma_k(f), \\ \Delta(F^{(2)}) &\equiv \frac{\Delta(f)}{4\alpha\tilde{a} - \mathfrak{A}^2} \pmod{\sigma_{n-1}(f) \cdot 2^{t+\partial_{n-2}(f)}}, \\ \omega_{k-2}(G^{(2)}) &= \omega_k(g), & \sigma_{k-2}(G^{(2)}) &= \sigma_k(g), \\ \Delta(G^{(2)}) &\equiv \frac{\Delta(g)}{4\alpha\tilde{a} - \mathfrak{A}^2} \pmod{\sigma_{n-1}(g) \cdot 2^{t+\partial_{n-2}(g)}} \end{aligned}$$

und

$$F^{(2)}(h; 2^t) = \frac{f(h; 2^t)}{\chi(h; 2^t)}, \quad G^{(2)}(h; 2^t) = \frac{g(h; 2^t)}{\chi(h; 2^t)}. \quad [\chi(h; 2^t) \neq 0]$$

Nehmen wir an, unser Satz sei bereits für Formen mit einer kleineren Anzahl von Variablen als n bewiesen, so finden wir $F^{(2)} \simeq G^{(2)} \pmod{2^t}$ und daraus auch $f_{(2)} \simeq g_{(2)} \pmod{2^t}$.

Inhaltsverzeichnis.

Kapitel	Seite
Begleitschreiben an die Académie des Sciences.	3
Inhaltsübersicht	4

Erster Teil.

Über die Reste quadratischer Formen.

I. Klassen quadratischer Formen. — Index, Invarianten und Ordnung einer Form.	9
II. Formenreste. — Die Invarianten $o(f)$ sind ganze Zahlen	13
III. Hauptformenreste und Hauptrepräsentanten für einen Modul N	21
IV. Bedingungen für die Existenz einer Ordnung O	26
V. Sätze über Hauptreste. — Grundformen für einen Modul N	32
VI. Formengruppen für einen Modul N . — Charaktere	37
VII. Über die Anzahl der Lösungen der Kongruenzen $f = \sum_{i=1}^n a_{ik} x_i x_k \equiv m \pmod{N}$	45
VIII. Bestimmung der Größen $f(h; q^t)$ in den einfachsten Fällen	50
IX. Charaktere der Hauptrepräsentanten und der Grundformen	58
X. Bedingungen für die Gültigkeit der Kongruenz $f \simeq g \pmod{q^t}$	70
XI. Genera von Formen. — Bedingungen für die Existenz eines Genus	71
XII. Adjungierte Formen. — Reziprozität zwischen den Ordnungen $\left(\begin{matrix} \sigma_1, \sigma_2, \dots, \sigma_{n-2}, \sigma_{n-1} \\ o_1, o_2, \dots, o_{n-2}, o_{n-1} \end{matrix} \right), I$ und $\left(\begin{matrix} \sigma_{n-1}, \sigma_{n-2}, \dots, \sigma_2, \sigma_1 \\ o_{n-1}, o_{n-2}, \dots, o_2, o_1 \end{matrix} \right), I$	80

Zweiter Teil.

Über die Darstellung ganzer Zahlen durch quadratische Formen.

XIII. Hilfssatz	83
XIV. Darstellung einer Form von ν Variablen durch eine Form von n Variablen $(\nu < n)$. — Äquivalente Darstellungen und Darstellungsgruppen	86
XV. Adjungierte Darstellungen und adjungierte Darstellungsgruppen	91
XVI. Darstellung von ganzen Zahlen durch Formen mit n Variablen	94
XVII. Darstellungen von Formen mit $n-1$ Variablen durch Formen mit n Variablen	95
XVIII. Index, Ordnung und Genus der durch eine Form von n Variablen dar- stellbaren Formen von $n-1$ Variablen	102
XIX. Über den Inbegriff der Darstellungen einer ganzen Zahl durch die ver- schiedenen Formen eines Genus G	113
XX. Maß eines positiven Genus. — Maß der Darstellungen einer ganzen Zahl durch die Formen eines positiven Genus.	116
XXI. Über die Anzahl der Darstellungen einer ganzen Zahl durch eine Summe von fünf Quadraten	117
XXII. Über die Bestimmung des Maßes einiger positiver Genera.	119
XXIII. Maß eines beliebigen Genus einer Ordnung $\left(\begin{matrix} 1 \\ o_h \end{matrix} \right) [o_h \equiv 1 \pmod{2}]$	134
Note über die Kongruenzen $f \simeq g \pmod{q^t}$	136

Verzeichnis der hauptsächlichsten vom Herausgeber übersetzten Stellen.*)

Das *Begleitschreiben* (S. 3—4) und die *Inhaltsübersicht* (S. 4—9) finden sich, diese in deutscher, jenes in französischer Sprache, nur im deutschen Manuskript, während sie in den *Mémoires des Savants étrangers* nicht mit zum Abdruck gelangt sind.

Erster Teil.

Überschrift (S. 9, Z. 21).

Kap. I. S. 9, Fußnote. S. 10, Z. 1; Z. 6; Z. 34. S. 12, Z. 28—29.

Kap. II. S. 13, Z. 16—22. S. 14, Z. 2—3; Z. 28—30; Fußnote. S. 15, Z. 1; Z. 21; Z. 32—33. S. 16, Z. 20—29; Z. 31—S. 17, Z. 28. S. 18, Z. 5—6; Z. 15—16; Z. 28—31. S. 19, Z. 19—20; Z. 30—33. S. 20, Z. 3.

Kap. III. S. 24, Z. 8; Z. 13; Z. 16; Z. 36. S. 25, Z. 25—28.

Kap. IV. S. 26, Z. 1; Z. 22—28; Z. 32—34. S. 27, Z. 5—6.

Kap. V. S. 33; Z. 9. S. 34, Z. 9—11; Z. 26—27. S. 35, Z. 17—18.

Kap. VI. S. 37, Z. 25—31. S. 39, Z. 8; Z. 28—S. 40, Z. 33. S. 41, Z. 5. S. 42, Z. 16—17. S. 43, Z. 22. S. 45, Z. 1; Z. 9—10.

Kap. VII. S. 45, Z. 17—18; Z. 24—28. S. 47, Fußnote. S. 48, Z. 11; Z. 15.

Kap. VIII. S. 50, Z. 12—13; Z. 19—20. S. 51, Z. 10—13; Z. 16—17. S. 53, Z. 20. S. 58, Z. 8—10.

Kap. IX. S. 58, Z. 17—23. S. 62, Z. 13—14. S. 65, Z. 5—10. S. 69, Z. 18; Z. 25.

Kap. X. S. 70, Z. 11—S. 71, Z. 4.

Kap. XI. S. 71, Z. 5—36. S. 72, Z. 12—14; Z. 22—23; Z. 27—28. S. 74, Z. 21. S. 77, Fußnote.

Kap. XII. S. 81, Z. 3—4; Z. 11; Z. 13—15. S. 82, Z. 23—24.

Zweiter Teil.

Überschrift (S. 83, Z. 9—10).

Kap. XIII. S. 83, Z. 24—S. 84, Z. 5; Z. 9—13; Z. 18; Z. 24—30. S. 86, Z. 7—8.

Kap. XIV. S. 86, Z. 24—27. S. 88, Z. 13—14; Z. 17—27. S. 89, Z. 1—S. 90, Z. 2; Z. 5—7; Z. 13—14.

Kap. XV. S. 91, Z. 8—10; Z. 25—S. 93, Z. 2; Z. 6—7.

Kap. XVI. S. 94, Z. 26—27.

Kap. XVII. S. 95, Z. 29—S. 96, Z. 9. S. 98, Z. 2; Z. 13—22; Z. 26—S. 99, Z. 8. S. 100, Z. 5—6; Z. 9—10; Z. 17—23. S. 102, Z. 10—12.

Kap. XVIII. S. 102, Z. 20—22; Z. 28—30. S. 103, Z. 4—14. S. 104, Z. 20; Z. 23—25; Z. 33. S. 106, Z. 8; Z. 18—24; Z. 32—33. S. 107, Z. 10; Z. 27—29. S. 108, Z. 6. S. 109, Z. 4. S. 110, Z. 28. S. 112; Z. 18—19; Z. 22.

Kap. XIX. S. 114, Z. 3—6; Z. 25—27. S. 115, Z. 18—22.

Kap. XX. S. 116, Z. 23—25.

Kap. XXI. S. 118, Z. 26—29; Z. 37.

Kap. XXII. S. 119, Z. 30. S. 120, Z. 15—16. S. 123, Z. 18—19. S. 124, Z. 20—21. S. 125, Z. 31. S. 129, Z. 23—28. S. 130, Z. 25. S. 131, Z. 6. S. 133, Z. 1—3; Z. 19—21; Z. 29. S. 134, Z. 3—4; Z. 15.

Kap. XXIII. S. 134, Z. 22; Z. 27—S. 135, Z. 5. S. 136, Z. 14—20.

Die *Note über die Kongruenzen* $f \not\equiv g \pmod{q^i}$ (S. 136—142) und das *Inhaltsverzeichnis* (S. 143) fehlen im deutschen Manuskript vollständig.

*) Vgl. die Vorbemerkung auf Seite 3. — S. bedeutet „Seite“, Z. „Zeile“. Bei der Abzählung der Zeilen sind die Symbole für Ordnungen und Unterdeterminanten

$$\begin{pmatrix} \sigma_1, \sigma_2, \dots, \sigma_{n-1} \\ \sigma_1, \sigma_2, \dots, \sigma_{n-1} \end{pmatrix}, I, \text{ bzw. } A \begin{pmatrix} i_1, i_2, \dots, i_h \\ k_1, k_2, \dots, k_h \end{pmatrix}, S \begin{pmatrix} l_1, l_2, \dots, l_h \\ i_1, i_2, \dots, i_h \end{pmatrix}$$

einzeilig, die als Zahlssysteme geschriebenen quadratischen Formen oder Substitutionen hingenom mit der Anzahl ihrer Druckzeilen in Ansatz gebracht.

II.

Sur la réduction des formes quadratiques positives quaternaires.*)

(Comptes rendus de l'Académie des Sciences, Paris 1883, t. 96, pp. 1205—1210).

(Note de M. Minkowski, présentée par M. Jordan.)

M. Charve a publié, aux Comptes rendus de 1881, une Note sur la réduction des formes quadratiques positives quaternaires. Je prends la liberté d'ajouter sur cet objet les observations suivantes, auxquelles je suis parvenu en étudiant les beaux Mémoires de M. Hermite dans le Journal de Crelle t. 40, 41 et 47. (Oeuvres, t. I, p. 94, p. 100, p. 164, p. 193, p. 200, p. 234.)

Les recherches sur la réduction des formes quadratiques positives s'appuient sur le théorème suivant:

I. Une forme quadratique positive $f = \sum_{h,k=1}^n a_{hk} x_h x_k$ à déterminant D n'obtient que pour un nombre fini de systèmes numériques (x_h) une valeur qui ne surpasse pas une quantité positive donnée E .

Démonstration. En appliquant à f la substitution au déterminant 1,

$$x_h = y_h, \quad x_k = y_k + \frac{\frac{\partial D}{\partial a_{hk}}}{\frac{\partial D}{\partial a_{hh}}} y_h \quad (k \neq h),$$

f sera transformée en

$$g = \frac{D}{\frac{\partial D}{\partial a_{hh}}} y_h^2 + g_h \quad (y_h = x_h),$$

où g_h représente une forme positive aux $n - 1$ variables y_k ($k \neq h$). Par conséquent, si la valeur de f ne doit pas être plus grande que la quantité E , nous obtenons un système d'inégalités

$$E \geq f = g \geq \frac{D}{\frac{\partial D}{\partial a_{hh}}} x_h^2,$$

auquel ne peut satisfaire qu'un nombre fini de nombres entiers x_h .

*) Der Text dieser Arbeit, wie er hier veröffentlicht wird, folgt einem Manuskript, das von Minkowski angefertigt wurde, nachdem die Arbeit bereits in den Comptes rendus erschienen war, und das weit korrekter ist, als die dort zum Abdruck gelangte ältere Fassung. (Anm. d. Herausg.)

Nous arrangeons toutes les formes positives en un certain ordre.

Nous disons qu'une forme positive $g = \sum_{h,k=1}^n b_{hk} x_h x_k$ est placée à côté d'une forme $f = \sum_{h,k=1}^n a_{hk} x_h x_k$, si les coefficients b_{hk} sont égaux aux coefficients a_{hk} , par contre *au-dessus* (ou *au-dessous*) de la forme f , si les coefficients b_{hk} et a_{hk} ne sont pas tous d'accord, et si le premier coefficient b_{hh} , qui n'est pas égal au coefficient correspondant a_{hh} , est plus grand (ou plus petit) que a_{hh} .

A l'aide du théorème I on peut facilement démontrer: 1^o que, parmi toutes les formes qui résultent d'une forme donnée f à l'aide des substitutions numériques au déterminant 1 et qui donnent la classe f , apparaissent certaines formes φ qui sont placées au-dessous de toutes les autres formes de cette classe; 2^o que le nombre de ces formes φ est ordinairement égal à 1, et que ce n'est qu'exceptionnellement qu'il devient plus grand que 1; 3^o que ces formes φ peuvent toujours être trouvées par un procédé fini.

Les formes $\varphi = \sum_{h,k=1}^n \alpha_{hk} \xi_h \xi_k$ sont ce que M. Hermite appelle des formes réduites. Les formes réduites φ de la même classe ont sûrement les mêmes coefficients α_{hk} .

II. Pour $n = 2, 3, 4$, une forme φ est réduite, et elle ne l'est que si elle satisfait aux inégalités

$$(I) \quad \alpha_{11} \leq \alpha_{22} \leq \dots \leq \alpha_{nn}$$

et à toutes les inégalités

$$(II) \quad \varphi(\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n) \geq \alpha_{hh},$$

dans lesquelles ε_h signifie une unité, et où les autres ε_k ($k \neq h$) ont ou les valeurs 0, ou +1, ou -1.

Démonstration. A) Les conditions (I) et (II) sont sûrement nécessaires pour que φ soit une forme réduite; car si l'on avait $\alpha_{hh} > \alpha_{kk}$ pour $h < k$, φ serait transformée par la substitution

$$S_I: \quad \xi_h = \eta_k, \quad \xi_k = \eta_h; \quad \xi_l = \eta_l \quad (h < k, l \neq h, k)$$

et si l'on avait $\varphi(\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n) < \alpha_{hh}$ ($\varepsilon_h = \pm 1, \varepsilon_k = 0, \pm 1$), par la substitution

$$S_{II}: \quad \xi_h = \varepsilon_h \eta_h, \quad \xi_k = \eta_k + \varepsilon_k \eta_h \quad (k \neq h)$$

en une forme qui serait placée au-dessous de φ ; par conséquent, φ ne pourrait pas être réduite. Des inégalités (II) résultent spécialement les conditions $\alpha_{hh} \pm 2\alpha_{hk} + \alpha_{kk} \geq \alpha_{hh}$, c'est-à-dire

$$(\alpha) \quad \alpha_{kk} \geq \pm 2\alpha_{hk}.$$

B) Les conditions (I) et (II) sont aussi *suffisantes* pour que φ soit réduite. Pour le démontrer, nous observons:

1° Que des inégalités (II) résultent toutes les autres inégalités

$$(m) \quad \varphi(m_1, m_2, \dots, m_n) \geq \alpha_{hh} \quad (m_h > 0),$$

dans lesquelles m_h signifie un nombre quelconque différent de zéro, et les autres m_k des nombres entiers tout à fait à volonté.

Nous désignons par ε_k le nombre 0, ou +1, ou -1, selon que m_k est égal à zéro, ou positif, ou négatif. Les quantités $\varepsilon_k m_k = \mu_k$ représentent les valeurs absolues des nombres m_k . Si les n quantités m_k , à l'exception du seul nombre m_h , sont égales à zéro, $\varphi(m_k)$ devient évidemment égal à $\alpha_{hh} m_h^2 \geq \alpha_{hh}$ et l'inégalité (m) aura lieu. Mais si des nombres m_k , au moins deux, sont différents de zéro, nous déterminons un indice t de la manière suivante. Nous cherchons les nombres m_τ , dont la valeur absolue est la plus petite sans être nulle, et nous posons, si parmi ces nombres m_τ se trouve aussi le nombre m_h , l'indice $t = h$, mais si ce n'est pas le cas, t égal à l'un quelconque des nombres τ . Si nous écrivons alors $s_k = \varepsilon_k \mu_t (k \neq t)$ et $s_t = 0$, les nombres s_k ne seront pas tous égaux à zéro, et nous obtenons l'identité:

$$\begin{aligned} \varphi(m_1, m_2, \dots, m_n) &= \varphi(m_k - s_k + s_k) = \varphi(m_k - s_k) + \sum_{i, k=1}^n \alpha_{ik} (m_i s_k + m_k s_i - s_i s_k) \\ &= \varphi(m_k - s_k) + 2m_t^2 \sum_{k \neq t} \alpha_{kt} \varepsilon_k \varepsilon_t + \sum_{i, k \neq t} \alpha_{ik} (m_i s_k + m_k s_i - s_i s_k), \end{aligned}$$

qui, à cause de la relation

$$2 \sum_{k \neq t} \alpha_{kt} \varepsilon_k \varepsilon_t = [\varphi(\varepsilon_k) - \alpha_{tt}] - \sum_{i, k \neq t} \alpha_{ik} \varepsilon_i \varepsilon_k,$$

prend la forme

$$\varphi(m_k) = \varphi(m_k - s_k) + m_t^2 [\varphi(\varepsilon_k) - \alpha_{tt}] + 2\mu_t \sum_{i \neq t} \varepsilon_i (\mu_i - \mu_t) \sum_{k \neq t} \alpha_{ik} \varepsilon_k.$$

Ici, par suite des inégalités (II) et (α), ni la quantité $[\varphi(\varepsilon_k) - \alpha_{tt}]$, ni les quantités

$$\sum_{k \neq t} \alpha_{ik} \varepsilon_i \varepsilon_k \left(= \alpha_{it} \varepsilon_i^2, \alpha_{it} \varepsilon_i^3 + \alpha_{ik} \varepsilon_i \varepsilon_k, \alpha_{it} \varepsilon_i^2 + \alpha_{ik} \varepsilon_i \varepsilon_k + \alpha_{ik'} \varepsilon_i \varepsilon_{k'} \right)$$

ne sont négatives; par conséquent, nous obtenons l'inégalité

$$\varphi(m_k) \geq \varphi(m_k - s_k),$$

pendant qu'on a en même temps

$$m_h - s_h > 0 \quad \text{et} \quad \sum m_k^2 > \sum (m_k - s_k)^2.$$

Si nous mettons maintenant $\sum m_k^2 = M$, et si nous supposons que le point 1° du théorème B) soit déjà prouvé pour tous les systèmes $(m_1, m_2, \dots, m_n)$, $m_h > 0$, pour lesquels $\sum m_k^2$ devient plus petit que M ,

il ressort évidemment $\varphi(m_k - s_k) \geq \alpha_{hh}$ et $\varphi(m_k) \geq \alpha_{hh}$. Sans doute l'inégalité (m) a lieu pour les systèmes (m_k) , pour lesquels on a $\sum m_k^2 \leq 1$, $m_h > 0$, c'est-à-dire pour le seul système $m_h = 1$, $m_k = 0$ ($k \neq h$). Ainsi, le point 1^o est parfaitement démontré.

2^o Admettons que pour la forme φ les inégalités (I) et (II) soient satisfaites et, par conséquent, aussi toutes les inégalités (m). Alors φ est, en effet, réduite.

Démonstration. Supposons d'abord que φ puisse être transformée par une substitution numérique $\xi_h = \sum_{k=1}^n r_h^k \eta_k$ à déterminant 1 en une forme $\psi = \sum_{h,k=1}^n \beta_{hk} \eta_h \eta_k$ qui soit placée au-dessous de φ .

Soit β_{ii} le premier des coefficients $\beta_{11}, \beta_{22}, \dots, \beta_{nn}$ de la forme ψ qui n'est pas égal au coefficient correspondant de la forme φ . On aura $\beta_{kk} = \alpha_{kk}$ ($k < i$) et $\beta_{ii} = \varphi(r_1^i, r_2^i, \dots, r_n^i) < \alpha_{ii}$.

Soit r_i^i la dernière des quantités $r_1^i, r_2^i, \dots, r_n^i$ qui est différente de zéro. On aura $\beta_{ii} = \varphi(\dots, r_i^i, \dots) \geq \alpha_{ii}$. Soit α_{tt} ($t \geq i$) la dernière des quantités $\alpha_{11}, \alpha_{22}, \dots, \alpha_{nn}$, qui a encore la valeur α_{ii} tandis que α_{hh} devient $> \alpha_{ii}$, si l'indice h est $> t$. Des relations $\alpha_{ii} > \beta_{ii}$, $\beta_{ii} \geq \alpha_{ii} = \alpha_{tt}$ et des inégalités (I) on conclut que le nombre t est $< i$; donc pour $k \leq t$ on aura $\beta_{kk} = \alpha_{kk} \leq \alpha_{tt}$. Par conséquent, toutes les quantités r_h^k , pour lesquelles on a $k \leq t$ et $h > t$, seront égales à zéro, puisque chacune de ces quantités, si elle différait de zéro, fournirait l'inégalité $\alpha_{kk} \geq \alpha_{hh}$, tandis que l'on a $\alpha_{kk} \leq \alpha_{tt}$, $\alpha_{hh} > \alpha_{tt}$. Dans le déterminant $|r_h^k|$ s'évanouissent donc toutes les $(t+1)(n-t)$ quantités

$$r_h^k (k = 1, 2, \dots, t; i; h = t+1, \dots, n)$$

qui sont les termes d'un système de $n-t$ séries horizontales et de $t+1$ séries verticales. Ce déterminant doit donc être égal à zéro, et l'on rencontre une contradiction.

De ce qui précède on déduit, à l'aide du théorème I, que l'on peut déterminer pour chaque forme positive f toutes les formes réduites de sa classe à l'aide d'un nombre fini de substitutions de la forme S_I ou S_{II} . (Toute classe f , pour laquelle aucune des inégalités (I) et (II) ne se change en une équation, a une seule forme réduite.)

Dans le cas $n = 4$, j'ai trouvé pour la limitation des coefficients des formes réduites des identités symétriques analogues à celles que Gauss a établies pour les formes ternaires (Gauss, Oeuvres complètes, t. II). Je reviendrai sur ces identités dans une autre occasion.

III.

Über positive quadratische Formen.

(Crelles Journal für die reine und angewandte Mathematik, Bd. 99, S. 1—9.)

Meine Abhandlung „Sur la théorie des formes quadratiques à coefficients entiers“*) schließt mit der Bestimmung des Maßes für einige Genera von positiven quadratischen Formen. Das *Maß* eines Genus ist eine Größe, welche erhalten wird, indem man sämtliche Klassen des Genus abzählt und dabei die einzelnen Klassen in entsprechenden Verhältnissen rechnet, als sie verschiedene Formen besitzen. Ich hatte mich seinerzeit auf die Betrachtung spezieller Genera beschränkt. Mittlerweile bin ich zu einfachen Ergebnissen für das Maß eines beliebigen Genus gelangt. Die schließlichen Formeln sind zu einem Teile bereits von H. J. Stephen Smith veröffentlicht worden**). Das Resultat gewinnt aber außerordentlich an Klarheit und erlangt eine weitergehende Bedeutung, indem man jene Definition des Genus zugrunde legt, von welcher ich in der erwähnten Arbeit ausgegangen bin, und zu welcher auch Herr Poincaré geführt worden ist***).

Ich setze Formen von fester Variablenzahl n und von nichtverschwindender Determinante voraus. Dann definiere ich:

A. Das *Genus* einer Form $f = \sum_1^n a_{ik} x_i x_k$ wird gebildet von allen den

Formen g , welche denselben Trägheitsindex wie f besitzen und welche mit f für jeden beliebigen Modul N kongruent sind.

Dabei heißt eine Form g *kongruent* mit f in bezug auf einen Modul N , wenn es möglich ist, aus f mit Hilfe einer linearen Substitution von einer Determinante $\equiv 1 \pmod{N}$ eine Form herzuleiten, in welcher sämtliche Koeffizienten für den Modul N dieselben Reste lassen wie die entsprechenden Koeffizienten von g .

Die Definition A. stellt an die Formen g nur scheinbar unendlich viele Anforderungen. In Wirklichkeit gilt der Satz:

*) Mémoires présentés à l'Académie des Sciences T. XXIX. Nr. 2. Unter dem Titel „Grundlagen für eine Theorie der quadratischen Formen mit ganzzahligen Koeffizienten“, diese Ges. Abhandlungen, Bd. I, S. 3—144.

**) Proceedings of the Royal Society of London. 1868. (Collected Papers, vol. I, p. 510.)

***) Comptes rendus de l'Académie des Sciences à Paris. 1882. I.

B. Eine Form g gehört dann und nur dann dem Genus einer Form f an, wenn sie denselben Index I und dieselbe Determinante Δ wie f besitzt und dazu mit f für den Modul 2Δ kongruent ist.

Immer dann, wenn diese Bedingungen erfüllt sind, wird es zugleich (nach Sätzen von Smith) möglich sein, die Form f in die Form g mittels solcher linearer Substitutionen von der Determinante 1 überzuführen, in denen die Koeffizienten rationale Zahlen mit einem zu 2Δ relativ primen Generalnenner sind.

Das Genus einer Form $f = \{a_{ik}\}$ ist eindeutig bestimmt, sobald man sein *vollständiges System von Invarianten* kennt. Dieses System umfaßt die folgenden Größen:

1. Den Index I der Form f .
2. Die größten positiven Teiler der sämtlichen Unterdeterminanten von 1, 2, ..., n Reihen des Systems $|a_{ik}|$. Diese Teiler mögen der Reihe nach mit $d_0, d_1, \dots, d_{n-1}$ bezeichnet werden, so daß sich insbesondere

$$\Delta = (-1)^I \cdot d_{n-1}$$

ergibt.

3. $n - 1$ Größen $\sigma_1, \sigma_2, \dots, \sigma_{n-1}$, welche die Werte 1 oder 2 haben. Und zwar ist ein σ_h gleich 1, wenn unter den symmetrischen h -reihigen Minoren von $|a_{ik}|$ sich solche vorfinden, die, vom Teiler d_{h-1} befreit, ungerade ausfallen; gleich 2, wenn das Entgegengesetzte eintritt.

4. Eine Reihe von Einheiten ± 1 , die Charaktere der Form f .

In Art. IV und Art. XI meiner am Anfange zitierten Arbeit [[diese Ges. Abhandlungen, Bd. I, S. 31—32 und S. 75—76]] sind alle Bedingungen zusammengestellt, denen die Invarianten eines (wirklich existierenden) Genus zu genügen haben. Insbesondere müssen die n Gleichungen

$$d_0 = o_0, d_1 = o_0^2 o_1, \dots, d_{n-1} = o_0^n o_1^{n-1} \dots o_{n-1}$$

stets zu n ganzen Zahlen $o_0, o_1, \dots, o_{n-1}$ führen.

Ich will hier nur von *positiven* Formen sprechen, also $I = 0$ voraussetzen. Jede positive Form f läßt eine endliche Anzahl von Transformationen von der Determinante 1 in sich zu. Diese Anzahl mag $t(f)$ genannt sein. Die Größe $t(f)$ ist konstant für alle Formen der Klasse f ; ihr reziproker Wert stellt das Maß der Klasse f vor.

Das Maß M eines positiven Genus f ($= \{a_{ik}\}$) ist nun definiert durch eine Summe:

$$\sum \frac{1}{t(\varphi)},$$

erstreckt über sämtliche verschiedenen Formenklassen φ , welche das betrachtete Genus aufweist*).

*) Eisenstein, Crelles Journal, Bd. 35. (Mathem. Abhandlungen, S. 177.)

Mit $f(N)$ will ich bezeichnen, wie viele für einen Modul N inkongruente Substitutionen von einer Determinante $\equiv 1 \pmod{N}$ man bilden kann, die, auf $f = \{a_{ik}\}$ angewandt, die sämtlichen Reste $a_{ik} \pmod{N}$ ungeändert lassen. Gemäß der Definition A. wird die Zahl $f(N)$ eine Invariante des Genus f vorstellen. Den reziproken Wert dieser Zahl wird man passend als das *Maß* des Genus f für den Modul N bezeichnen können.

Diese speziellen Maße $\frac{1}{f(N)}$ finde ich als die wesentlichen Faktoren des Maßes M .

In betreff der Zahlen $f(N)$ gelten folgende Sätze:

1. Wenn N sich aus mehreren Primzahlpotenzen zusammensetzt, $N = \prod q^t$, so ist $f(N) = \prod f(q^t)$.

Für Potenzen einer festen Primzahl q findet man:

2. Für alle Potenzen q^t , welche die höchste in

$$2 \cdot \prod_{h=0}^{n-1} o_h$$

enthaltene Potenz von q überschreiten, fällt die Größe

$$\frac{f(q^t)}{q^{\frac{n(n-1)}{2}t}}$$

konstant aus. Man setze diese Größe gleich

$$q^U \cdot f\{q\},$$

wo q^U die höchste in $\prod_{h=0}^{n-1} \left(o_h^{\frac{(n-h)(n-h+1)}{2}-1}\right)$ enthaltene Potenz von q sei.

Alsdann wird $f\{q\}$ im wesentlichen nur von quadratischen Charakteren des Genus f abhängen. (Ferner wird für das zu f adjungierte Genus

$$F = \sum_1^n \frac{\partial \Delta}{\partial a_{ik}} x_i x_k$$

die Gleichung $F\{q\} = f\{q\}$ gelten.) Die Größe $f\{q\}$ bildet so den hauptsächlichen Faktor aller Größen $f(q^t)$.

Wenn nun

$$D = \prod_{h=1}^{n-1} o_h^{h(n-h)},$$

so erhalte ich für das Maß des Genus f einfach:

$$(1) \quad M = c \cdot \sqrt{D} \cdot \frac{1}{f\{2\}} \cdot \frac{1}{f\{3\}} \cdot \frac{1}{f\{5\}} \cdots \frac{1}{f\{q\}} \cdots,$$

wo c die Konstante

$$2 \frac{\Gamma\left(\frac{1}{2}\right) \cdot \Gamma\left(\frac{2}{2}\right) \cdots \Gamma\left(\frac{n}{2}\right)}{\left\{\Gamma\left(\frac{1}{2}\right)\right\}^{\frac{n(n+1)}{2}}} \quad \left(\Gamma\left(\frac{1}{2}\right) = \sqrt{\pi}, \text{ usw.}\right)$$

bedeutet, und wo das Produkt der Reihe nach über alle Primzahlen $q = 2, 3, 5, 7, \dots$ auszudehnen ist.

Als Beispiel betrachte man ein positives Genus von (eigentlich) primitiven binären Formen von einer Determinante $\Delta \equiv 1 \pmod{4}$. Man hat hier $n = 2$, $\sigma_1 = 1$, $\sigma_1 = D = \Delta$. Für jede ungerade Primzahl q , die nicht in Δ aufgeht, ergibt sich

$$f\{q\} = 1 - \left(\frac{-\Delta}{q}\right) \frac{1}{q};$$

dagegen für jede Primzahl q aus 2Δ :

$$f\{q\} = 2.$$

Bedeutet also μ die Anzahl aller (ungeraden) Primzahlen von Δ , so kommt

$$M = \frac{1}{2^\mu} \cdot \frac{\sqrt{\Delta}}{\pi} \sum \left(\frac{-\Delta}{m}\right) \frac{1}{m},$$

wo m alle positiven und zu 2Δ relativ primen Zahlen zu durchlaufen hat. Ist $\Delta > 1$, so besitzt jede Klasse unseres Genus das Maß $\frac{1}{2}$. Die Größe M wird dann gleich der halben Klassenanzahl des betrachteten Genus. Man gewinnt so ein Resultat, welches mit den Dirichletschen Formeln übereinstimmt.

Ohne mich bei weiteren Beispielen aufzuhalten, will ich nur zeigen, daß der Ausdruck (1) stets einen endlichen Wert besitzt.

Für jede ungerade Primzahl q , die nicht in D aufgeht, findet man, wenn n gerade:

$$f\{q\} = \left(1 - \frac{1}{q^2}\right) \left(1 - \frac{1}{q^4}\right) \cdots \left(1 - \frac{1}{q^{n-2}}\right) \cdot \left\{1 - \left(\frac{(-1)^{\frac{n}{2}} D}{q}\right) \frac{1}{q^{\frac{n}{2}}}\right\},$$

und wenn n ungerade:

$$f\{q\} = \left(1 - \frac{1}{q^2}\right) \left(1 - \frac{1}{q^4}\right) \cdots \left(1 - \frac{1}{q^{n-1}}\right).$$

Für jede solche Primzahl fällt also die Größe

$$\frac{\left(1 - \frac{1}{q^2}\right) \left(1 - \frac{1}{q^4}\right) \cdots \left(1 - \frac{1}{q^{\frac{n-1}{2}}}\right)}{f\{q\}} = E_q$$

besonders einfach aus. Nun kann man leicht in (1) diese Größe als allgemeines Produktglied einführen. Man braucht nur mit der Identität

$$1 = S_2 S_4 \cdots S_{2\left[\frac{n-1}{2}\right]} \cdot \prod_q \left(1 - \frac{1}{q^2}\right) \left(1 - \frac{1}{q^4}\right) \cdots \left(1 - \frac{1}{q^{\frac{n-1}{2}}}\right)$$

zu multiplizieren, wo S_{2k} die Summe $1 + \frac{1}{2^{2k}} + \frac{1}{3^{2k}} + \dots$ bezeichnet, deren Wert in bekannter Weise durch die k^{te} Bernoullische Zahl, B_k , ausgedrückt ist. Setzt man, wenn n gerade:

$$c = \left(\frac{1}{2}\right)^{\frac{n-4}{2}} \cdot B_1 B_2 \dots B_{\frac{n-2}{2}},$$

und wenn n ungerade:

$$c = \left(\frac{1}{2}\right)^{\frac{n-3}{2}} \cdot B_1 B_2 \dots B_{\frac{n-1}{2}} \cdot \frac{1}{1 \cdot 2 \dots \frac{n-1}{2}},$$

läßt man ferner ∂ die sämtlichen Primzahlen von $2D$ durchlaufen, so kommt, wenn $n \equiv 0 \pmod{2}$:

$$(2) \quad M = c \cdot \frac{\sqrt{D}}{\pi^{\frac{n}{2}}} \cdot \prod_{\partial} E_{\partial} \cdot \sum \left(\frac{(-1)^{\frac{n}{2}} D}{m} \right) \frac{1}{m^{\frac{n}{2}}},$$

(wo die Summation alle positiven und zu $2D$ relativ primen Zahlen m betrifft), und wenn $n \equiv 1 \pmod{2}$, so kommt:

$$(2) \quad M = c \cdot \sqrt{D} \cdot \prod_{\partial} E_{\partial}.$$

Diese Ausdrücke sind denen ähnlich, welche Smith angegeben hat, und man kann wohl aus der erwähnten Note von Smith die Werte der Größen E_{∂} für ungerade Primzahlen ∂ , sowie in einigen speziellen Fällen auch für die Primzahl $\partial = 2$ entnehmen. Doch tritt dort nicht die hier entwickelte einfache Bedeutung dieser Größen zutage, und diese gerade ist es, welche erkennen läßt, daß ähnliche Verhältnisse auch für allgemeinere Formen bestehen. Vor allem fällt auf, daß der Ausdruck (1) nichts enthält, was an *positive* quadratische Formen gebunden ist. In der Tat besitzt dieser Ausdruck auch für alle indefiniten Genera (von nicht zerlegbaren Formen) einen endlichen Wert. Indem man nach der Bedeutung dieses Wertes forscht, gelangt man zu einer interessanten Erklärung des Maßbegriffes für indefinite Formen.

Ich will noch einige Sätze in betreff der Reduktion der positiven Formen mit beliebigen reellen Koeffizienten hinzufügen. Ist diese es ja, welche die Mittel gibt, um in jedem einzelnen Falle die Klassen eines Genus zu sondern und deren Maß zu berechnen. Ich wende (in etwas veränderter Form) diejenige Reduktionsmethode an, welche Herr Hermite im 40. Bande des Crelleschen Journals S. 302 (Oeuvres, T. I, p. 149) aufgestellt hat.

„In einer gegebenen Klasse von positiven quadratischen Formen sollen diejenigen Formen

$$f = \sum_1^n a_{ik} x_i x_k$$

reduzierte heißen, welche vor allen anderen Formen den kleinsten Wert der Verbindung

$$R = a_{11} + \delta \cdot a_{22} + \delta^2 \cdot a_{33} + \dots + \delta^{n-1} \cdot a_{nn}$$

ergeben, wo δ eine positive unendlich kleine Größe bezeichnet.“

Gemäß dieser Definition werden alle reduzierten Formen einer Klasse in den n Koeffizienten

$$(a_{11}, a_{22}, \dots, a_{nn})$$

übereinstimmen. Der Koeffizient a_{11} insbesondere wird das *Minimum* der Klasse f vorstellen.

Die charakteristischen Bedingungen solcher reduzierten Formen rühren für $n = 2$ von Lagrange und für $n = 3$ von Seeber her. Für den Fall $n = 4$ habe ich einen diesbezüglichen Satz in den Comptes rendus vom April 1883 [[diese Ges. Abhandlungen, Bd. I, S. 145—148]] mitgeteilt. (Der Beweis ist dort, namentlich am Schluß, durch eine Menge Druckfehler sehr entstellt.) Ein weiterer Satz existiert nun für den Fall $n = 5$. Ich fasse das Resultat für alle vier Fälle $n = 2, 3, 4, 5$ zusammen.

„In diesen Fällen ist eine Form

$$f = \sum_1^n a_{ik} x_i x_k$$

immer und nur dann eine positive und reduzierte Form, wenn sie einer Reihe von Ungleichungen

$$(I) \quad f(m_1, m_2, \dots, m_n) \geq a_{ii} \quad (i = 1, 2, \dots, n)$$

genügt, wo die m_h , je nach den vier Fällen, den folgenden Tabellen gemäß zu wählen sind:

$n = 2,$

m_i	$\pm m_{i'}$
1	1

$n = 3,$

m_i	$\pm m_{i'}$	$\pm m_{i''}$
1	1	0
	1	1

$n = 4,$

m_i	$\pm m_{i'}$	$\pm m_{i''}$	$\pm m_{i'''}$
1	1	0	0
	1	1	0
	1	1	1

$n = 5,$

m_i	$\pm m_{i'}$	$\pm m_{i''}$	$\pm m_{i'''}$	$\pm m_{i^{IV}}$
1	1	0	0	0
	1	1	0	0
	1	1	1	0
	1	1	1	1
	1	1	1	2

und wenn ferner

$$(II) \quad a_{11} \leq a_{22} \leq \dots \leq a_{nn}$$

ist.“

In allen diesen Fällen wird also eine Hermitesche reduzierte Form durch eine endliche Anzahl einzelner linearer Ungleichungen definiert.

Genügt eine Form allen Bedingungen (I), so ist sie entweder selbst reduziert oder läßt sich doch durch eine oder mehrere Substitutionen von der Art

$$x_i = y_k, \quad x_k = -y_i$$

in eine reduzierte Form überführen. Jedenfalls stimmen ihre n Hauptkoeffizienten, abgesehen von der Reihenfolge, mit den n Koeffizienten a_{ii} einer ihr äquivalenten reduzierten Form überein. Nun erkennt man leicht, daß den Bedingungen (I) alle solchen Formen genügen müssen, welche in ihrer Klasse den kleinsten Wert der Verbindung

$$\Re_1 = a_{11} + a_{22} + \dots + a_{nn}$$

oder der Verbindung

$$\Re_n = a_{11} a_{22} \dots a_{nn}$$

oder gewisser ähnlicher Verbindungen $\Re$ ergeben.

In den Fällen $n = 2, 3, 4, 5$ liefern sonach die Hermiteschen reduzierten Formen einer Klasse für jede einzelne der Größen $a_{11} + a_{22} + \dots + a_{nn}$, $a_{11}a_{22} + a_{11}a_{33} + \dots$, $a_{11}a_{22} \dots a_{nn}$, ... den kleinsten Wert, welchen diese Größe für Formen der betrachteten Klasse überhaupt annehmen kann. Umgekehrt, wenn eine Form von n ($= 2, 3, 4, 5$) Variablen für irgendeine der Verbindungen $\Re$ einen kleinsten Wert ergibt, so geht diese Form bei geeigneter Permutation ihrer Koeffizientenreihen in eine Hermitesche reduzierte Form über; sie liefert also auch für alle anderen Verbindungen $\Re$ einen kleinsten Wert.

Diese Sätze, welche für $n = 2$ und $n = 3$ interessante Eigenschaften ebener und räumlicher Gitter ausdrücken, scheinen um so bemerkenswerter, als sie nicht mehr für größere Werte des n gelten.

Im allgemeinen ist eine reduzierte Form um so eigentümlicher, in je mehr von den Ungleichungen (I) und (II) das Gleichheitszeichen statt hat. Es gibt nun positive reduzierte Formen f , für welche $\frac{n(n+1)}{2} - 1$ linear unabhängige von den Ungleichungen (I) und (II) in Gleichungen übergehen, durch die dann die Verhältnisse der $\frac{n(n+1)}{2}$ Größen a_{ik} vollständig und zwar in rationaler Weise bestimmt sind. Solche reduzierten Formen mögen *Grenzformen* heißen. Unter allen Grenzformen, denen dieselben Verhältnisse der Koeffizienten zukommen, wird es offenbar eine geben, deren Koeffizienten ganze Zahlen ohne einen gemeinsamen Teiler sind. Diese heiße eine *primitive Grenzform*.

„Da die Anzahl der Ungleichungen (I) und (II) eine beschränkte ist, so existiert (für $n = 2, 3, 4, 5$) immer nur eine beschränkte Anzahl von primitiven Grenzformen.“

Um dieselben wirklich darzustellen, will ich mit $(a)_\nu$ diejenige quadratische Form von ν Variablen bezeichnen, deren ν Hauptkoeffizienten sämtlich gleich a , und deren übrige Koeffizienten sämtlich gleich 1 sind. Die Determinante dieser Form ist $(a-1)^{\nu-1} \cdot (a-1+\nu)$. Ich finde dann folgende primitive Grenzformen:

Wenn $n = 2$ ist, die Form $\varphi = (2)_2 = 2(x_1^2 + x_1x_2 + x_2^2)$ und deren Nebenreduzierte $2(x_1^2 - x_1x_2 + x_2^2)$.

Wenn $n = 3$ ist, die Form $\varphi = (2)_3$ und die Nebenreduzierten dieser Form.

Wenn $n = 4$ ist, die Form $\varphi = (2)_4$ und deren Nebenreduzierten; ferner die Form ψ , welche durch die Substitution

$$x_h = y_h + y_4, \quad x_4 = 2y_4 \quad (h = 1, 2, 3)$$

in

$$(2)_1 + (2)_1 + (2)_1 + (2)_1 = 2(y_1^2 + y_2^2 + y_3^2 + y_4^2)$$

übergeht, und die Nebenreduzierten dieser Form ψ .

Wenn $n = 5$ ist, die Form $\varphi = (2)_5$ und deren Nebenreduzierten; weiter die Form ψ , welche durch die Substitution

$$x_h = y_h + y_5, \quad x_5 = -y_4 + y_5 \quad (h = 1, 2, 3, 4)$$

in $(2)_1 + (2)_1 + (2)_3$ übergeht, und deren Nebenreduzierten; endlich die Form χ , welche durch

$$x_h = y_h + (-1)^h y_5, \quad x_5 = 2y_5 \quad (h = 1, 2, 3, 4)$$

in $(4)_5$ übergeht, und die Nebenreduzierten von χ .

Ohne Mühe wird man erkennen, daß die Klassen dieser Grenzformen identisch sind mit den „*formes extrêmes*“, welche die Herren Korkine und Zolotareff in ihren interessanten Aufsätzen im 6. und 11. Bande der Mathematischen Annalen eingeführt haben. Es sind darunter positive Formen verstanden, welche die Eigenschaft besitzen, daß ihr Minimum abnimmt, so oft den Koeffizienten unendlich kleine Variationen beigelegt werden, welche die Determinante der Form nicht vergrößern. Die nahe Beziehung zu diesen „Formen mit größtem Minimum“ ist für die Reduktionsmethode von Herrn Hermite charakteristisch.

Wiesbaden, den 31. Januar 1885.

IV.

Untersuchungen über quadratische Formen.

Bestimmung der Anzahl verschiedener Formen, welche ein gegebenes Genus enthält.

(Inauguraldissertation (Königsberg 1885); Acta Mathematica, Band 7, S. 201—258.)

Dem Andenken meines teuren Vaters.

In meiner Arbeit «*Sur la théorie des formes quadratiques à coefficients entiers*»*) habe ich den Begriff des Genus lediglich aus dem Begriffe der Formenkongruenz hergeleitet. Ein solches Verfahren erwies sich bereits dort als äußerst vorteilhaft. Seine Berechtigung wird vielleicht noch schärfer durch die folgenden Entwicklungen hervortreten. Ich werde mich hier mit jenen Zahlen beschäftigen, welche im Falle allgemeiner Genera dieselbe Rolle spielen, wie im Falle binärer Genera die Klassenanzahlen. Gewisse Sätze über Formenkongruenzen werden zunächst erraten lassen, in welcher Gestalt sich die Ausdrücke jener Zahlen darbieten müssen. Unter Anwendung Dirichletscher Prinzipien soll alsdann gezeigt werden, daß die erratenen Formeln in Wirklichkeit richtig sind.

Einleitung.

1. Ein Genus und seine Formenanzahl.

Unter den *Resten* einer quadratischen Form $f = \sum_1^n a_{ik} x_i x_k$ in bezug auf einen Modul N verstehen wir alle diejenigen Formen, welche aus f hervorgehen, indem die Koeffizienten a_{ik} in beliebiger Weise um Vielfache des Moduls N geändert werden.

*) Mémoires présentés par divers savants à l'Académie des Sciences de l'Institut de France. Tome XXIX, N°. 2 (1884). Ich zitiere diese Arbeit im folgenden kurz mit *F. Q.* — Den auf *F. Q.* bezüglichen Zitaten fügen wir stets in [[]] den entsprechenden Hinweis auf die hier (Bd. I, S. 3—144) veröffentlichte deutsche Ausgabe hinzu. (Anm. d. Herausg.)

Wir nennen zwei Formen f und g (von derselben Variablenzahl n) *kongruent* in bezug auf einen Modul N , $f \cong g \pmod{N}$, wenn es lineare Substitutionen von einer Determinante $\equiv 1 \pmod{N}$ gibt, durch welche die Reste der einen Form für den Modul N in Reste der anderen Form für den Modul N übergehen.

Wir betrachten ausschließlich Formen von nichtverschwindender Determinante.

Es kann der Fall eintreten, daß zwei Formen f und g für einen jeden beliebigen Modul kongruent sind. Solches findet offenbar immer statt, wenn f und g derselben Klasse äquivalenter Formen angehören, d. i. durch ganzzahlige Substitutionen von der Determinante 1 ineinander übergehen. Es ereignet sich überhaupt dann und nur dann, wenn f und g dieselbe Determinante Δ besitzen und in bezug auf den Modul 2Δ kongruent sind.

Immer und nur dann, wenn die Formen f und g diesen Bedingungen genügen und dazu einen gleichen Trägheitsindex liefern, wird es möglich sein, die eine dieser Formen in die andere durch solche linearen Substitutionen von der Determinante 1 überzuführen, in welchen die Koeffizienten rationale Zahlen mit einem zu 2Δ relativ primen Generalnenner sind.*)

Alle Formen, welche denselben Trägheitsindex I haben wie eine gegebene Form, und welche mit dieser Form für einen jeden beliebigen Modul kongruent sind, fassen wir in ein *Genus* zusammen. Da die Formen eines Genus eine feste Determinante besitzen, so können sie, nach bekannten Sätzen, nur in eine endliche Anzahl verschiedener Formenklassen zerfallen.

Es ist aber im allgemeinen nicht sowohl die Anzahl dieser *Klassen*, welche sich durch einfache Formeln ausdrücken läßt, als vielmehr die Anzahl der in einem Genus enthaltenen *Formen*. Zwar ist die letztere Anzahl stets eine unendliche, denn schon jede einzelne Formenklasse besitzt unendlich viel Formen. Doch zeigt es sich bald, daß für ein jedes Genus eine gewisse, positiv unendliche Größe Ω existiert, welche nur von den einfachsten Invarianten des Genus abhängt, und zu welcher alle die Formenanzahlen der einzelnen Klassen des Genus in endlichen Verhältnissen stehen. Dieses Ω bestimmt dann auch den Grad, in welchem die Formenanzahl des gesamten Genus unendlich wird. Fällt die Formenanzahl in einer oder in mehreren Klassen des Genus gleich $M \cdot \Omega$ aus, so bezeichnen wir die positive endliche Größe M als das *Maß* der betreffenden Klassen.

*) Henry I. Stephen Smith, Philosophical Transactions, CLVII, 1867 (On the Orders and Genera of Ternary Quadratic Forms, art. 12) und: Proceedings of the Royal Society of London, XVI, 1868 (On the Orders and Genera of Quadratic Forms containing more than three Indeterminates, p. 202). (Collected Papers, vol. I, p. 480 und p. 516.)

Am einfachsten gestaltet sich der Begriff des Maßes für definite Formen, also in den Fällen $I=0$ und $I=n$, mit welchen wir uns später hauptsächlich beschäftigen werden. Man weiß, daß eine definite quadratische Form f immer nur eine endliche Anzahl von Transformationen von der Determinante 1 in sich zuläßt. Sei $t(f)$ diese Anzahl. Nennen wir ferner Ω_0 die Anzahl aller möglichen linearen ganzzahligen Substitutionen S von n Reihen und von der Determinante 1. Würden wir alle diese S auf die Form f anwenden, so müßten wir zu einer jeden Form der Klasse f gelangen, und zwar zu einer jeden genau $t(f)$ -mal. Die Anzahl der verschiedenen Formen der Klasse f wird daher $\frac{1}{t(f)} \cdot \Omega_0$ betragen. Wählen wir demnach, was in der Tat geschehen soll, im Falle definiter Formen die erwähnte Größe $\Omega = \Omega_0$, so können wir als das Maß einer definiten Klasse f die Größe $\frac{1}{t(f)}$ erklären. Das Maß für die Formenanzahl eines definiten Genus wird dann $\sum \frac{1}{t(f)}$ sein, wo die Summe über alle verschiedenen Klassen f des Genus zu erstrecken ist.

In solchen speziellen Fällen, wo die Größe $t(f)$ konstant ausfällt für alle Formen eines Genus, ergibt die vorstehende Summe einfach den $t(f)^{\text{ten}}$ Teil der Klassenanzahl des Genus. Derartige Fälle sind Regel bei binären Formen. Man hat da gewöhnlich $t(f) = 2$. Nur für die Klassen $f = d(x^2 + y^2)$ wird $t(f) = 4$, und für die Klassen $f = 2d(x^2 + xy + y^2)$ wird $t(f) = 6$.

Ich bemerke noch beiläufig, daß die Größe Ω_0 sich mit der $(n^2 - 1)^{\text{ten}}$ Potenz der Anzahl ω aller möglichen ganzen Zahlen von $-\infty$ bis $+\infty$ vergleichen läßt. Es gilt nämlich in gewissem Sinne die Beziehung:

$$\Omega_0 = \frac{\omega^{n^2-1}}{S_2 S_3 \dots S_n},$$

wo

$$S_k = 1 + \frac{1}{2^k} + \frac{1}{3^k} + \frac{1}{4^k} + \dots \quad (k=2, 3, \dots, n)$$

Eine Deutung des Maßbegriffes für indefinite Formen soll den Gegenstand einer späteren Arbeit bilden. Ich erwähne nur, daß als Maß einer primitiven binären Form $ax^2 + 2bxy + cy^2$ von einer negativen, nicht quadratischen Determinante $ac - b^2 = -D < 0$ am passendsten der Ausdruck $\frac{\pi}{\log \left(\frac{T + U\sqrt{D}}{\sigma} \right)}$ festgesetzt wird, wo $\sigma (= 1, 2)$ den größten Teiler

der Koeffizienten $a, 2b, c$ bedeutet, und wo T und U die zwei kleinsten positiven Zahlen sind, welche der Gleichung $T^2 - DU^2 = \sigma^2$ Genüge leisten.

2. Das vollständige System von Invarianten eines Genus.

Um ein Genus eindeutig zu charakterisieren, bedarf es nicht durchaus der Kenntniss einer seiner Formen. Es genügt bereits, wenn man imstande ist, sein *vollständiges System von Invarianten* anzugeben. Wir verstehen unter dieser Bezeichnung eine Reihe von Größen, welche sich als arithmetische Funktionen einer repräsentierenden Form f des Genus darstellen lassen, und welche die doppelte Eigenschaft haben, einmal: ungeändert zu bleiben, so oft die Form f durch eine andere Form *desselben* Genus ersetzt wird, dann aber: in ihrer Gesamtheit sich stets zu ändern, sobald für f irgendeine Form eines *anderen* Genus genommen wird.

Ein vollständiges System von Invarianten eines Genus f umfaßt die folgenden Größen:

1. den Trägheitsindex I der Form f . Derselbe sagt aus, wie viele Quadrate mit dem Vorzeichen Minus auftreten, sobald f durch irgendeine reelle Substitution in eine Summe von n , zum Teil positiven, zum Teil negativen Quadraten transformiert wird;

2. die n Invarianten d_0 ; $o_1, o_2, \dots, o_{n-1}$, und die $n-1$ Invarianten $\sigma_1, \sigma_2, \dots, \sigma_{n-1}$ der Form f .

Die Invariante d_0 stellt den positiven größten gemeinsamen Teiler der Koeffizienten a_{ik} von f vor.

Zu den Invarianten o_h gelangt man in folgender Weise. Man bezeichne mit $d_1, d_2, \dots, d_{n-1}$ die positiven größten gemeinsamen Teiler aller 2-, 3-, $\dots$, n -reihigen Unterdeterminanten von $|a_{ik}|$, so daß insbesondere $\Delta = (-1)^I \cdot d_{n-1}$ kommt. Aus den Zahlen d_h entstehen die Invarianten o_h vermittels der Gleichungen:

$$\frac{d_1}{d_0^2} = o_1, \quad \frac{d_2}{d_0^3} = o_1^2 o_2, \quad \dots, \quad \frac{d_{n-1}}{d_0^n} = o_1^{n-1} o_2^{n-2} \dots o_{n-1},$$

oder

$$o_h = \frac{d_h d_{h-2}}{d_{h-1}^2}.$$

Die Invarianten σ sind Größen von den Werten 1 oder 2. Man hat $\sigma_h = 1$, wenn unter den symmetrischen h -reihigen Minoren von f sich solche vorfinden, die, von dem Faktor d_{h-1} befreit, ungerade ausfallen; dagegen $\sigma_h = 2$, wenn diese Minoren alle durch $2d_{h-1}$ aufgehen.

Die Größen o_h sind stets ganze Zahlen, und die Größen σ_h müssen den folgenden Bedingungen genügen:

I. Die Zahlen $\sigma_{h-1} o_h \sigma_{h+1}$ und σ_h dürfen nicht zugleich durch 2 teilbar sein.

II. Die Quotienten $\frac{\sigma_{h-1} o_h}{\sigma_{h+1}}$ und $\frac{o_h \sigma_{h+1}}{\sigma_{h-1}}$ müssen ganz sein.

Wir führen noch die Invarianten ein: $\sigma_0 = 1$, $\sigma_n = 1$ und $o_0 = 0$, $o_n = 0$.

3. die *Charaktere* der Form f . Dieses sind eine Reihe von Einheiten ± 1 , welche in Gestalt Legendrescher Symbole auftreten. Um sie möglichst einfach darzustellen, trifft man am besten folgende Voraussetzung über die repräsentierende Form f : Die aus den ersten $h(=1, 2, \dots, n-1)$ Reihen von f gebildeten symmetrischen Minoren sollen Werte $\sigma_h d_{h-1} \varphi_h$ ergeben von solcher Art, daß ein jedes φ_h relativ prim zu $2o_1 o_2 \dots o_{n-1}$ und zu φ_{h-1} und φ_{h+1} ausfällt. Dabei hat man sich noch $\varphi_0 = 1$ und $\varphi_n = (-1)^I$ zu denken. Nach *F. Q.*, p. 83 [[S. 72]], lassen sich in jeder Klasse des Genus Formen f von dieser Eigentümlichkeit finden; man nennt sie *charakteristische Formen*.

Es sei allgemein ϵ_h das Vorzeichen von φ_h ; ferner mag I_h angeben, wie viele von den Einheiten $\frac{\epsilon_1}{\epsilon_0}, \frac{\epsilon_2}{\epsilon_1}, \dots, \frac{\epsilon_h}{\epsilon_{h-1}}$ negativ ausfallen, so daß insbesondere $I_0 = 0$ und $I_n = I$ zu nehmen ist. Für eine charakteristische Form f erweisen sich folgende Einheiten C als Charaktere:

1) wenn $\sigma_{h-1} o_h \sigma_{h+1}$ durch eine ungerade Primzahl p teilbar ist, die Einheit

$$\left(\frac{\varphi_h}{p}\right);$$

2) wenn $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{4}$ ist, die Einheiten

$$(-1)^{\frac{\varphi_h - 1}{2}}, \quad \left(\frac{\varphi_{h-1}}{\epsilon_h \varphi_h}\right) (-1)^{\frac{I_h(I_h - 1)}{2}}, \quad \left(\frac{\varphi_{h+1}}{\epsilon_h \varphi_h}\right) (-1)^{\frac{I_h(I_h + 1)}{2}};$$

3) wenn $\sigma_{h-1} o_h \sigma_{h+1} \equiv 0 \pmod{8}$ ist, die Einheit

$$\left(\frac{2}{\varphi_h}\right).$$

Diese Einheiten C bilden zusammen mit den Größen I, d_0, o_h, σ_h wirklich ein vollständiges System von Invarianten für das betrachtete Genus, so daß kein zweites Genus da sein kann, welches zu eben diesen Invarianten führt.

Die Einheiten C müssen alle Bedingungen erfüllen, welche sich für sie aus den quadratischen Kongruenzen

$$(1) \quad -\sigma_{h-1} o_h \sigma_{h+1} \varphi_{h-1} \varphi_{h+1} \equiv X_h^2 \pmod{\sigma_h^2 \varphi_h}$$

erschließen lassen. Man findet diese Bedingungen in *F. Q.*, pp. 87—88 [[S. 75—76]] zusammengestellt.

Es gilt der Satz:

(A) Wenn die Invarianten I, d_0, o, σ und die Charaktere C in beliebiger Weise festgesetzt werden, doch so, daß sie den Bedingungen I., II. genügen, ferner allen Bedingungen, welche aus den Kongruenzen (1) folgen, so existiert wirklich ein Genus, welches zu diesen Invarianten und Charakteren Veranlassung gibt.

Ein Beweis dieses Satzes ist *F. Q.*, pp. 89—90 [[S. 76—77]], mit Hilfe des bekannten Theorems über die arithmetischen Progressionen geführt. Ich habe dort diesen Hilfssatz $(n-1)$ -mal hintereinander angewandt zur sukzessiven Auffindung von $n-1$ Zahlen $\varphi_1, \varphi_2, \dots, \varphi_{n-1}$ für eine charakteristische Form des gewünschten Genus. Man überzeugt sich aber leicht, daß es immer genügt, ein einziges Mal von diesem Satze Gebrauch zu machen, nämlich um allein φ_{n-1} als Primzahl zu bestimmen; und man sieht dann ferner, daß in den Fällen $n \geq 3$ dieser Hilfssatz sich ganz entbehren läßt, wenn nur der Satz (A) bereits für $n=2$ bewiesen ist.

Alle Formengenera, welche dieselben Werte der Größen I, d_0, o_h, σ_h besitzen, werden in eine *Ordnung*

$$d_0, \quad \left(\begin{matrix} \sigma_1, \sigma_2, \dots, \sigma_{n-1} \\ o_1, o_2, \dots, o_{n-1} \end{matrix} \right), \quad I$$

zusammengefaßt.

Wir beschäftigen uns meist mit Formen f , deren d_0 gleich 1 ist, und die man *primitive* Formen nennt. Im Falle die Größe d_0 einer Form f relativ prim zu einer Zahl N ist, heißt diese Form primitiv in bezug auf N .

Erster Teil.

3. Die Anzahl der Reste eines Genus in bezug auf einen Modul N .

Wir wollen zunächst untersuchen, wieviel verschiedene Formenreste ein gegebenes Genus in bezug auf einen gegebenen Modul N liefert. In der Anzahl dieser Formenreste finden wir einen Divisor der Formenanzahl des Genus, und wir werden so alle wesentlichen Faktoren kennen lernen, aus welchen sich die Ausdrücke für die Formenanzahl zusammensetzen.

Das Genus möge durch eine beliebige seiner Formen, f , repräsentiert werden. Es ist dann klar, daß wir die Reste des Genus nach dem Modul N nur unter denjenigen Formen suchen dürfen, welche $\equiv f \pmod{N}$ sind. Man gelangt zu diesen Formen, indem man auf f ein vollständiges System von lauter inkongruenten Substitutionen T von einer Determinante $\equiv 1 \pmod{N}$ anwendet. Ein solches System wird leicht bestimmt. Man braucht beispielsweise nur von den N^n inkongruenten Substitutionen

$$x_i \equiv \sum_k t_i^k y_k \pmod{N}, \quad t_i^k = 1, 2, \dots, N \quad (i, k = 1, 2, \dots, n)$$

alle diejenigen fortzulassen, in welchen die Determinante nicht $\equiv 1 \pmod{N}$ ausfällt.

In dem Systeme der T mögen sich im ganzen $\mathfrak{N}$ Substitutionen finden lassen, — etwa die folgenden: $T_1, T_2, \dots, T_{\mathfrak{N}}$ —, durch welche f in $\mathfrak{N}$ verschiedene Reste: $g_1, g_2, \dots, g_{\mathfrak{N}} \pmod{N}$ übergehe. Das gegebene

Genus enthält dann sicher nicht mehr als diese $\mathfrak{N}$ Reste. Wir behaupten aber, daß diese Reste wirklich alle dem gegebenen Genus, ja schon der (ganz beliebigen) Klasse f eigen sind.

In der Tat, zu jeder Substitution

$$T = \begin{vmatrix} x_1, & \dots & \\ x_2, & \dots & \\ \dots & \dots & \\ x_n, & \dots & \end{vmatrix} \equiv 1 \pmod{N}$$

läßt sich immer eine nach dem Modul N kongruente Substitution S von einer Determinante 1 bestimmen. Im Falle $n=1$ leuchtet dieses unmittelbar ein. Wenn $n > 1$ ist, so bedienen wir uns zum Beweise eines Schlusses von $n-1$ auf n . Offenbar muß der größte Teiler der n Zahlen x_i zu N relativ prim sein. Man kann daher n Zahlen $\xi_i \equiv x_i \pmod{N}$ finden, deren größter Teiler gleich 1 ist. Alsdann läßt sich bekanntlich eine Substitution

$$S_0 = \begin{vmatrix} \xi_1, & \dots & \\ \xi_2, & \dots & \\ \dots & \dots & \\ \xi_n, & \dots & \end{vmatrix}$$

von der Determinante 1 bilden (*F. Q.*, p. 98 [[S. 83]]). Das Produkt $S_0^{-1} \cdot T$ gewinnt jetzt die Form

$$U \equiv \begin{vmatrix} 1, & U_1, & \dots, & U_{n-1} \\ 0, & u_1^1, & \dots, & u_1^{n-1} \\ \dots & \dots & \dots & \dots \\ 0, & u_{n-1}^1, & \dots, & u_{n-1}^{n-1} \end{vmatrix} \equiv 1 \pmod{N}.$$

Ist unser Lemma bereits für den Fall $n-1$ erwiesen, so können wir anstelle der u_i^k solche kongruente Zahlen einführen, daß $|u_i^k|$ in 1 übergeht. Ferner setzen wir anstelle der ersten Vertikalreihe von U einfach die Zahlen 1, 0, ..., 0. Auf diese Weise wird U in eine kongruente Substitution V von der Determinante 1 übergehen, und das Produkt $S_0 \cdot V = S$ erscheint als eine mit T kongruente Substitution von der Determinante 1.

Wir können so zu allen Substitutionen $T_1, T_2, \dots, T_{\mathfrak{N}}$ kongruente Substitutionen $S_1, S_2, \dots, S_{\mathfrak{N}}$ von der Determinante 1 bilden. Durch diese muß sich f in äquivalente Formen mit den Resten $g_1, g_2, \dots, g_{\mathfrak{N}}$ verwandeln, was zu beweisen war.

Es ist nun leicht plausibel zu machen, daß die Formenanzahl der Klasse f ein Vielfaches der Zahl $\mathfrak{N}$ wird. Man denke sich zu dem Behufe in der Klasse f alle verschiedenen Formen gekennzeichnet, welche

nach dem Modul N den Rest f lassen, und für jede dieser Formen je eine Substitution S notiert, durch welche dieselbe aus f entsteht. Durch Anwendung aller $S \cdot S_1, S \cdot S_2, \dots, S \cdot S_{\mathfrak{N}}$ müssen dann aus f die sämtlichen Formen der Klasse f hervorgehen, und zwar eine jede ein einziges Mal. Die Zahl $\mathfrak{N}$ teilt somit wirklich in gewissem Sinne die (unendliche) Formenanzahl der Klasse f , und da diese Klasse eine beliebige ist, auch die Formenanzahl des gesamten Genus, wie am Anfange ausgesprochen war.

Um einen Ausdruck für die Zahl $\mathfrak{N}$ zu gewinnen, verfährt man folgendermaßen. Es sei $\psi_n(N)$ die Anzahl der Individuen eines vollständigen Systemes von inkongruenten Substitutionen T von einer Determinante $\equiv 1 \pmod{N}$. Alle diejenigen T , welche auf den Rest $f \pmod{N}$ ohne Wirkung bleiben, nenne man T , und es sei $f(N)$ die Anzahl der verschiedenen T . Die Substitutionen $\mathsf{T} \cdot T_1, \mathsf{T} \cdot T_2, \dots, \mathsf{T} \cdot T_{\mathfrak{N}}$ müssen das gesamte System der T erschöpfen, und man erhält so:

$$\mathfrak{N} = \frac{\psi_n(N)}{f(N)}.$$

In welcher Weise die Größe $\psi_n(N)$ gefunden wird, ist bekannt*). Damit ein $T \equiv 1 \pmod{N}$ ausfalle, ist zunächst erforderlich, daß die n Zahlen $x_1, x_2, \dots, x_n$ der ersten Vertikalreihe ohne gemeinsamen Teiler mit N gewählt seien. Eine solche Wahl kann auf $N^n \cdot (N)_n$ Arten geschehen, wenn $(N)_n$ das über alle verschiedenen Primzahlen q von N ausgedehnte Produkt $\prod \left(1 - \frac{1}{q^n}\right)$ bedeutet. Man sieht leicht, daß zu jedem, ohne Teiler mit N gewählten Systeme x_i mindestens ein T gehört. Alle möglichen T mit der ersten Vertikalreihe x_i folgen dann durch Zusammensetzung dieses einen mit den verschiedenen Substitutionen $U \equiv 1 \pmod{N}$. In U unterliegen die $n-1$ Zahlen U_i gar keiner Beschränkung. Es lassen sich also im ganzen $N^{n-1} \cdot \psi_{n-1}(N)$ inkongruente U bilden, und man erlangt die Beziehung:

$$\psi_n(N) = N^{2n-1} \cdot (N)_n \cdot \psi_{n-1}(N).$$

Da nun $\psi_1(N) = 1$ ist, so entsteht

$$\psi_n(N) = N^{n^2-1} \cdot (N)_2 \cdot (N)_3 \cdots (N)_n.$$

Es handelt sich also wesentlich um die Bestimmung der Größen $f(N)$. Dabei genügt es, den Fall zu untersuchen, wo N eine Primzahlpotenz q^t ist.

Denn setzt N sich aus mehreren Primzahlpotenzen q^t zusammen, $N = \prod q^t$, so hat man $f(N) = \prod f(q^t)$.

*) Jordan, Traité des substitutions, 120—124.

So oft nämlich eine Substitution $T \equiv 1 \pmod{N}$ ist, ergibt sich dieselbe auch $\equiv 1$ nach jedem der Moduln q^t , und ändert sie den Rest $f \pmod{N}$ nicht, so ändert sie auch keinen der Reste $f \pmod{q^t}$. Liegt andererseits für einen jeden Modul q^t ein T_q vor, von einer Determinante $\equiv 1 \pmod{q^t}$, welches auf $f \pmod{q^t}$ ohne Wirkung bleibt, so wird die Substitution $T \pmod{N}$, welche allen Kongruenzen

$$T \equiv T_q \pmod{q^t}$$

genügt, $\equiv 1 \pmod{N}$ ausfallen, und den Rest $f \pmod{N}$ nicht ändern. So geht die behauptete Relation hervor.

4. Hilfssätze zur Bestimmung der Zahlen $f(q^t)$.

Wir brauchen in betreff der Größen $f(q^t)$ nur den Fall zu betrachten, wo f primitiv in bezug auf q ist, also die Koeffizienten a_{ik} von f nicht sämtlich den Teiler q haben. Denn sei etwa $f = q^d \cdot g$ und $d > 0$. So lange man $d \geq t > 0$ hat, bleibt der Rest $f \pmod{q^t}$ bei jeder Substitution ungeändert, und man erhält $f(q^t) = \psi_n(q^t)$. Wenn aber $d < t$ ist, so gilt die Relation

$$(2) \quad f(q^t) = q^{(n^2-1)d} \cdot g(q^{t-d}).$$

In der Tat, eine jede Substitution $T \equiv 1 \pmod{q^t}$, welche auf $f \pmod{q^t}$ ohne Wirkung ist, ändert auch $g \pmod{q^{t-d}}$ nicht; und ebenso wird ein jedes $\mathfrak{T} \equiv 1 \pmod{q^{t-d}}$, welches auf $g \pmod{q^{t-d}}$ ohne Wirkung bleibt, auch $f \pmod{q^t}$ nicht ändern. Aus einem jeden $\mathfrak{T}$ lassen sich aber $q^{(n^2-1)d}$, nach dem Modul q^t verschiedene Substitutionen T herleiten. Denn in einem $\mathfrak{T}$ ist immer mindestens ein Koeffizient c da, für welchen $\frac{\partial \mathfrak{T}}{\partial c}$ zu q relativ prim ausfällt. Wir können nun, um eine Substitution T zu gewinnen, erst jeden der $n^2 - 1$ übrigen Koeffizientenreste $\pmod{q^{t-d}}$ von $\mathfrak{T}$ durch q^d verschiedene Reste $\pmod{q^t}$ ersetzen. Der Rest von c für den Modul q^t folgt hernach eindeutig aus der Bedingung, daß die veränderte Substitution eine Determinante $\equiv 1 \pmod{q^t}$ ergebe.

Insofern es uns um Faktoren für die Formenanzahl eines Genus zu tun ist, reicht die Betrachtung solcher Moduln q^t aus, welche gewisse Grenzen $q^{G(q)}$ überschreiten. Denn ist $t \geq t - d > 0$, so beweist man leicht, daß die Zahl $\psi_n(q^{t-d}) : f(q^{t-d})$ einen Divisor der Zahl $\mathfrak{N} = \frac{\psi_n(q^t)}{f(q^t)}$ vorstellt. Beachtet man die oben gegebenen Werte von $\psi_n(q^t)$, so läuft dieses darauf hinaus, daß die Zahl $f(q^t)$ in der Zahl $q^{(n^2-1)d} \cdot f(q^{t-d}) [t - d > 0]$ aufgeht.

Für eine mit q primitive Form f machen wir von folgenden Bezeichnungen Gebrauch. Die höchsten in den Invarianten $o_1, o_2, \dots, o_{n-1}$ ent-

haltenen Potenzen von q sollen die Exponenten besitzen: $\omega_1, \omega_2, \dots, \omega_{n-1}$, und es sei allgemein

$$v_h = \omega_1 + \omega_2 + \dots + \omega_h, \quad \partial_h = h\omega_1 + (h-1)\omega_2 + \dots + \omega_h.$$

Ferner nehmen wir an, von den $n-1$ nicht negativen Zahlen ω_h seien im ganzen $\lambda-1$ größer als Null, nämlich die folgenden:

$$(\vartheta_0 = 0) \qquad \omega_{\vartheta_1}, \omega_{\vartheta_2}, \dots, \omega_{\vartheta_{\lambda-1}} \qquad (\vartheta_{\lambda} = n)$$

und wir bilden die Gleichungen

$$\vartheta_1 = \kappa_1, \quad \vartheta_2 = \kappa_1 + \kappa_2, \quad \dots, \quad \vartheta_{\lambda} = \kappa_1 + \kappa_2 + \dots + \kappa_{\lambda},$$

d. i.

$$\kappa_k = \vartheta_k - \vartheta_{k-1}.$$

Der Quotient aus der Determinante Δ von f und der Potenz $q^{\partial_{n-1}}$ mag für einen Moment Δ_0 heißen. Wir werden weiterhin voraussetzen, daß unsere Moduln q^t , wenn q einer ungeraden Primzahl p gleich ist, die Potenz $p^G = p^{v_{n-1}}$, und wenn $q = 2$ ist, die Potenz $2^G = 2^{1+v_{n-1}}$ überschreiten. Die Folge davon wird sein, daß ein jeder Rest $f \pmod{q^t}$ uns, wenn $q = p$ ist, den Wert der Einheit $\left(\frac{\Delta_0}{p}\right)$, und wenn $q = 2$ ist, den Wert der Einheit $(-1)^{\frac{\Delta_0-1}{2}}$, sowie im Falle $\sigma_{n-1} = 2$ auch den Wert der Einheit $\left(\frac{2}{\Delta_0}\right)$ liefert. Denn es gilt der Satz:

Genügt eine Form g schon der Kongruenz

$$g \cong f \pmod{q^{G+1}}$$

und soll noch $g \cong f \pmod{q^t}$ sein, so ist notwendig und hinreichend, daß, wenn $q = p$, die Beziehung

$$\Delta(g) \equiv \Delta(f) \pmod{p^{t+\partial_{n-2}}},$$

und wenn $q = 2$, die Beziehung

$$\Delta(g) \equiv \Delta(f) \pmod{\sigma_{n-1} \cdot 2^{t+\partial_{n-2}}}$$

bestehe.

Ich erwähne noch einige, zum Teil bekannte Sätze über Kongruenzen, welche bald ihre Anwendung finden werden.

(B) Ist eine Form f und eine Zahl α prim in bezug auf eine ungerade Primzahl p , und hat die Kongruenz

$$f(\xi_i) \equiv \alpha \pmod{p}$$

$A \cdot p^{n-1}$ Lösungen, so besitzt die Kongruenz

$$(3) \qquad f(\xi_i) \equiv \alpha \pmod{p^t} \qquad (t > 1)$$

$A \cdot p^{(n-1)t}$ Lösungen.

Denn genügt ein System ξ_i der Kongruenz

$$(4) \qquad f(\xi_i) \equiv \alpha \pmod{p^{t-1}},$$

und setzt man $x_i \equiv \xi_i + p^{t-1}u_i \pmod{p^t}$, so kommt, da $t > 1$ sein soll,

$$f(x_i) \equiv f(\xi_i) + p^{t-1} \cdot \sum u_i \frac{\partial f}{\partial \xi_i} \pmod{p^t}.$$

Nun können die Zahlen $\frac{\partial f}{\partial \xi_i}$ nicht sämtlich durch p teilbar sein, da man $\frac{1}{2} \sum \xi_i \frac{\partial f}{\partial \xi_i} \equiv \alpha \pmod{p}$ hat. Also ergibt die Kongruenz

$$\frac{\alpha - f(\xi_i)}{p^{t-1}} \equiv \sum u_i \frac{\partial f}{\partial \xi_i} \pmod{p}$$

p^{n-1} Lösungen $u_i \pmod{p}$, und eine jede Lösung von (4) liefert p^{n-1} Lösungen von (3), woraus unmittelbar unser Satz folgt.

Jetzt sei $f = \sum a_{ik} x_i x_k$ eine in bezug auf 2 primitive Form, und α eine ungerade Zahl. Wir unterscheiden zwei Fälle, je nachdem die Invariante σ_1 von f den Wert 1 oder 2 hat, die Zahlen a_{ii} also zum Teil ungerade oder sämtlich gerade ausfallen.

(C) Ist $\sigma_1 = 1$, und besitzt die Kongruenz

$$f(\xi_i) \equiv \alpha \pmod{8}$$

$A \cdot 2^{3(n-1)}$ Lösungen, so liefert die Kongruenz

$$(3) \quad f(\xi_i) \equiv \alpha \pmod{2^t} \quad (t > 3)$$

$A \cdot 2^{(n-1)t}$ Lösungen.

In der Tat, es genügen der Kongruenz

$$(4) \quad f(\xi_i) \equiv \alpha \pmod{2^{t-1}}$$

zusammen mit einem Systeme $\xi_i \pmod{2^{t-1}}$ immer alle die 2^n Systeme $x_i \pmod{2^{t-1}}$, für welche $x_i \equiv \xi_i \pmod{2^{t-2}}$ ist. Denn jedes dieser Systeme läßt sich in die Form $x_i \equiv \xi_i + 2^{t-2} \delta_i \pmod{2^{t-1}}$ setzen, wo die n Größen δ_i entweder 0 oder 1 bedeuten; man hat also wirklich:

$$f(x_i) \equiv f(\xi_i) + 2^{t-1} \cdot \sum \delta_i \frac{1}{2} \frac{\partial f}{\partial \xi_i} \equiv \alpha \pmod{2^{t-1}}.$$

Bildet man nun mit Hilfe einer Lösung $\xi_i \pmod{2^{t-2}}$ von (4) ein System $x_i \equiv \xi_i + 2^{t-2} u_i \pmod{2^{t-1}}$, so kommt

$$f(x_i) \equiv f(\xi_i) + 2^{t-1} \cdot \sum u_i \frac{1}{2} \frac{\partial f}{\partial \xi_i} \pmod{2^t},$$

und die Kongruenz

$$\frac{\alpha - f(\xi_i)}{2^{t-1}} \equiv \sum u_i \frac{1}{2} \frac{\partial f}{\partial \xi_i} \pmod{2}$$

ergibt 2^{n-1} Lösungen $u_i \pmod{2}$. So führt eine jede Lösung $\xi_i \pmod{2^{t-2}}$ von (4) zu 2^{n-1} Lösungen $\xi_i \pmod{2^{t-1}}$ von (3), woraus die Richtigkeit unserer Behauptung erhellt.

Es ist auch klar, daß unser Satz gültig bleibt, wenn wir alle solchen Lösungen ξ_i ausschließen, die zugleich gewissen gegebenen linearen Kongruenzen nach dem Modul 2 genügen.

(D) Ist zweitens $\sigma_1 = 2$, und hat die Kongruenz

$$f(\xi_i) \equiv 2\alpha \pmod{4}$$

$2A \cdot 2^{2(n-1)}$ Lösungen, in welchen die n Zahlen $\frac{1}{2} \frac{\partial f}{\partial \xi_i}$ nicht sämtlich gerade sind, so liefert die Kongruenz

$$(3) \quad f(\xi_i) \equiv 2\alpha \pmod{2^t} \quad (t > 2)$$

$2A \cdot 2^{(n-1)t}$ Lösungen, bei welchen die n Zahlen $\frac{1}{2} \frac{\partial f}{\partial \xi_i}$ nicht sämtlich gerade sind.

Denn setzt man $\frac{1}{2}f = \varphi$, so gruppieren sich je 2^n Lösungen $\xi_i \pmod{2^t}$ von (3) zu je einer Lösung $\xi_i \pmod{2^{t-1}}$ von $\varphi(\xi_i) \equiv \alpha \pmod{2^{t-1}}$. Unser Satz geht so in einen analogen Satz in betreff des Ausdruckes φ über, welcher dann ähnlich bewiesen wird wie der Satz (B).

Wir schreiten nun zur Bestimmung der Größen $f(q^t)^*$. Dabei werden wir uns hauptsächlich auf diejenigen Resultate stützen, welche in der Note zu meiner am Anfange zitierten Arbeit enthalten sind (*F. Q.*, pp. 169—178 [[S. 136—143]]).

5. Der Fall einer ungeraden Primzahl.

Wir betrachten zunächst den einfacheren Fall, wo q gleich einer ungeraden Primzahl p ist. Die Form f sei primitiv in bezug auf p . Der Modul p^t möge die Potenz p^{n-1} überschreiten.

Da wir f durch jeden nach dem Modul p^t kongruenten Rest ersetzen dürfen, so können wir annehmen (*F. Q.*, p. 7 [[S. 14]]), f habe den Typus

$$f \equiv \alpha \xi^2 + F \pmod{p^t},$$

wo α zu p prim ist und F einen Rest von $n-1$ Variablen vorstellt. Die Koeffizienten von F müssen den Faktor p^{ω_1} enthalten, und setzen wir

$$F \equiv p^{\omega_1} f^{(1)} \pmod{p^t},$$

so fällt der Rest $f^{(1)}$ primitiv in bezug auf p aus, und die $n-2$ Invarianten $p^{\omega_h^{(1)}}$, welche diesem Reste angehören, erfüllen die Gleichungen:

$$\omega_{h-1}^{(1)} = \omega_h. \quad (h = 2, 3, \dots, n-1)$$

Wir denken uns die $f(p^t)$ verschiedenen Substitutionen T von einer Determinante $\equiv 1 \pmod{p^t}$ aufgestellt, welche den Rest $f \pmod{p^t}$ in sich selbst überführen. In jeder dieser Substitutionen muß die erste Vertikalreihe aus n Zahlen $\xi_i \pmod{p^t}$ bestehen, welche

$$f(\xi_i) \equiv \alpha \pmod{p^t}$$

ergeben. Der vorstehenden Kongruenz mögen $A \cdot p^{(n-1)t}$ verschiedene Systeme $\xi_i \pmod{p^t}$ Genüge leisten. Die Betrachtungen aus *F. Q.*, p. 170

*) In einem ausgezeichneten Falle, nämlich für die Form $f = x_1^2 + x_2^2 + \dots + x_n^2$, sind die Zahlen $f(q)$ von Herrn Jordan gegeben (*Traité des substitutions*, 201—214, *Ordre du groupe orthogonal*).

[[S. 137]], lassen erkennen, daß jedem dieser Systeme ξ_i wirklich Substitutionen T zukommen, welche den Rest f in sich selbst transformieren.

Es fragt sich, wie viele verschiedene T können aus einem bestimmten Systeme ξ_i hervorgehen. Ist T_0 eine erste dieser Substitutionen, so wird jedes überhaupt vorhandene T mit der ersten Vertikalreihe ξ_i in ganz bestimmter Weise zusammengesetzt sein als Produkt aus T_0 und aus zwei Substitutionen U und $\mathfrak{T}$ von der Form

$$U \equiv \begin{vmatrix} 1, & U_1, & \dots, & U_{n-1} \\ 0, & 1, & \dots, & 0 \\ \dots & \dots & \dots & \dots \\ 0, & 0, & \dots, & 1 \end{vmatrix}, \quad \mathfrak{T} \equiv \begin{vmatrix} 1, & 0, & \dots, & 0 \\ 0, & t_1^1, & \dots, & t_1^{n-1} \\ \dots & \dots & \dots & \dots \\ 0, & t_{n-1}^1, & \dots, & t_{n-1}^{n-1} \end{vmatrix} \pmod{p^t}.$$

Soll nun T den Rest f in sich selbst überführen, so ist nötig, daß auch $U \cdot \mathfrak{T}$ diesen Rest in sich selbst transformiere. Hierzu wieder ist erforderlich, daß die Relationen $U_h \equiv 0 \pmod{p^t}$ gelten, und daß die Substitution $\mathfrak{T}$ auf den Rest $F \pmod{p^t}$ ohne Wirkung bleibe. Die Anzahl der verschiedenen T mit der ersten Vertikalreihe ξ_i , welche $f \pmod{p^t}$ nicht ändern, wird daher gleich der Anzahl der verschiedenen $\mathfrak{T}$ sein, welche auf $F \pmod{p^t}$ ohne Wirkung sind, also gleich $F(p^t)$. Zieht man noch die Formel (2) in Betracht, so kommt schließlich

$$f(p^t) = p^{(n-1)t + \omega_1[(n-1)^2 - 1]} \cdot A \cdot f^{(1)}(p^{t - \omega_1})$$

Wir setzen nun allgemein:

$$f(p^t) = p^{\frac{n(n-1)}{2}t + \sum_{h=1}^{n-1} \omega_h \left(\frac{(n-h)(n-h+1)}{2} - 1 \right)} \cdot f\{p\}. \quad \left(t > \sum_{h=1}^{n-1} \omega_h \right)$$

Für den Rest $f^{(1)}$ bilden wir eine entsprechende Größe $f^{(1)}\{p\}$. Die vorstehende Relation verwandelt sich alsdann in:

$$(5) \quad f\{p\} = A \cdot f^{(1)}\{p\},$$

und diese Formel bleibt auch für $n = 1$ gültig, falls nur für eine Form F von Null Variablen $F\{p\} = \frac{1}{2}$ genommen wird.

Es handelt sich jetzt um die Bestimmung der Größe A . Nach dem Satze (B) muß $A \cdot p^{n-1}$ die Anzahl aller Lösungen von $f(\xi_i) \equiv \alpha \pmod{p}$ ausdrücken. Bedeutet α den Index der ersten von den Zahlen $\omega_1, \omega_2, \dots, \omega_{n-1}, \omega_n (= -\infty)$, welche nicht verschwindet, so können wir f von dem Typus voraussetzen:

$$\begin{aligned} f &\equiv \Phi + p^{\omega_\alpha} \cdot f^{(\alpha)} \\ \Phi &\equiv \alpha_1 \xi_1^2 + \alpha_2 \xi_2^2 + \dots + \alpha_\alpha \xi_\alpha^2 \end{aligned} \pmod{p^t}, \quad (\alpha_1 = \alpha)$$

wo ein jedes $\alpha_1, \alpha_2, \dots, \alpha_\alpha$ und ebenso der Rest $f^{(\alpha)}$ primitiv in bezug auf p ist (*F. Q.*, p. 19 [[S. 22]]). Wir erhalten dann $f \equiv \Phi \pmod{p}$, und

$A \cdot p^{x-1}$ muß die Anzahl der Lösungen von $\Phi \equiv \alpha \pmod{p}$ geben. Nun läßt sich diese letztere Anzahl nach bekannten Sätzen herleiten.*) Man gelangt so zu folgenden Beziehungen:

1. wenn $x \equiv 0 \pmod{2}$,

$$A = \left(1 - \theta \cdot p^{-\frac{x}{2}}\right); \quad \theta = \left(\frac{(-1)^{\frac{x}{2}} \cdot \alpha_1 \alpha_2 \dots \alpha_x}{p}\right);$$

2. wenn $x \equiv 1 \pmod{2}$,

$$A = \left(1 + \theta^1 \cdot p^{-\frac{x-1}{2}}\right), \quad \theta^1 = \left(\frac{(-1)^{\frac{x-1}{2}} \cdot \alpha_2 \alpha_3 \dots \alpha_x}{p}\right).$$

Im Falle $x = 1$ hat man sich $\theta^1 = 1$ zu denken.

Wir können weiter die Größe $f\{p\}$ auf $f^{(x)}\{p\}$ zurückführen. Man findet:

$$f\{p\} = \mathfrak{A} \cdot f^{(x)}\{p\},$$

wenn $\mathfrak{A}$ den folgenden Ausdruck bedeutet: im Falle $x \equiv 0 \pmod{2}$,

$$\mathfrak{A} = 2 \left(1 - \frac{1}{p^2}\right) \left(1 - \frac{1}{p^4}\right) \dots \left(1 - \frac{1}{p^{x-2}}\right) \cdot \left(1 - \frac{\theta}{p^{\frac{x}{2}}}\right);$$

und im Falle $x \equiv 1 \pmod{2}$,

$$\mathfrak{A} = 2 \left(1 - \frac{1}{p^2}\right) \left(1 - \frac{1}{p^4}\right) \dots \left(1 - \frac{1}{p^{x-1}}\right).$$

Für ein $x = 1$ stimmt diese Gleichung unmittelbar mit (5) überein, während sie für ein $x > 1$ leicht aus (5) mit Hilfe eines Schlusses von $x - 1$ auf x hervorgeht. Man braucht nur zu beachten, daß dem Reste $f^{(1)}$ in derselben Weise die Zahl $x - 1$ angehört wie dem Reste f die Zahl x .

Für alle ungeraden Primzahlen p , welche nicht in der Determinante Δ der Form f aufgehen, wird $x = n$, also $\alpha_1 \alpha_2 \dots \alpha_x \equiv \Delta \pmod{p}$, während $f^{(x)}$ sich als ein Rest von Null Variablen erweist. Für alle diese Primzahlen p kommt daher:

1. wenn $n \equiv 0 \pmod{2}$,

$$f(p^t) = p^{\frac{n(n-1)}{2}t} \cdot \left(1 - \frac{1}{p^2}\right) \left(1 - \frac{1}{p^4}\right) \dots \left(1 - \frac{1}{p^{n-2}}\right) \cdot \left\{1 - \left(\frac{(-1)^{\frac{n}{2}} \Delta}{p}\right) \frac{1}{p^{\frac{n}{2}}}\right\};$$

2. wenn $n \equiv 1 \pmod{2}$,

$$f(p^t) = p^{\frac{n(n-1)}{2}t} \cdot \left(1 - \frac{1}{p^2}\right) \left(1 - \frac{1}{p^4}\right) \dots \left(1 - \frac{1}{p^{n-1}}\right).$$

*) Lebesgue, Recherches sur les nombres, § 5 (Journal de Liouville, T. 2, pp. 266—275). — C. Jordan, Comptes rendus, 1866, I, pp. 687—690; Traité des substitutions, 197—200; Journal de Liouville, 2^{me} série, T. 17, p. 372. — F. Q. artt. VII—VIII [[S. 45—58]].

Um die Größen $f\{p\}$ vollständig darzustellen, wollen wir endlich f als *Hauptrest* für den Modul p^t voraussetzen. Falls die Bezeichnungen aus 4. gelten, so heißt dieses, f soll den Typus haben:

$$f \equiv \{ \Phi_1 + p^{\omega_{\mathfrak{I}_1}} [\Phi_2 + p^{\omega_{\mathfrak{I}_2}} [\Phi_3 + \dots + p^{\omega_{\mathfrak{I}_{\lambda-1}}} (\Phi_{\lambda}) \dots]] \} \pmod{p^t},$$

$$\Phi_k \equiv \alpha_1^{(k)} \xi_1^{(k)} \xi_1^{(k)} + \dots + \alpha_{\kappa_k}^{(k)} \xi_{\kappa_k}^{(k)} \xi_{\kappa_k}^{(k)} \pmod{p^{t-\omega_{\mathfrak{I}_k-1}}},$$

wo die $\alpha_h^{(k)}$ sämtlich zu p prim sind (F. Q., p. 19 [[S. 22]]).

Für einen Hauptrest f wird die Größe $f\{p\}$ gleich einem Produkte aus der Potenz $2^{\lambda-1}$ und aus λ Faktoren $\mathfrak{A}_k$, welche den λ einzelnen Resten Φ_k entsprechen und folgende Bedeutung haben:

$$\mathfrak{A}_k = \left(1 - \frac{1}{p^2}\right) \left(1 - \frac{1}{p^4}\right) \dots \left(1 - \frac{1}{p^{\frac{1}{2} \left\lfloor \frac{\kappa_k}{2} \right\rfloor}}\right) \cdot a_k,$$

wo $\left\lfloor \frac{\kappa_k}{2} \right\rfloor$ die größte in $\frac{\kappa_k}{2}$ enthaltene ganze Zahl vorstellt, und $a_k = 1$ ist im Falle $\kappa_k \equiv 1 \pmod{2}$, dagegen:

$$a_k = \left(1 + \theta_k \cdot p^{-\frac{\kappa_k}{2}}\right)^{-1}, \quad \theta_k = \left(\frac{(-1)^{\frac{\kappa_k}{2}} \cdot \Delta(\Phi_k)}{p}\right),$$

wenn $\kappa_k \equiv 0 \pmod{2}$ ist.

Die Richtigkeit dieser Formeln ergibt sich sofort mit Hilfe eines Schlusses von $\lambda - 1$ auf λ . Man bemerkt nämlich, daß dem Reste $f^{(x)}(\kappa = \kappa_1)$ in derselben Weise die Zahl $\lambda - 1$ zukommt wie dem Reste f die Zahl λ .

6. Der Fall der Primzahl 2.

Wir behandeln jetzt den Fall $q = 2$, welcher auf etwas umständlicherem Wege zu gleich einfachen Resultaten führt. Die Form f sei primitiv in bezug auf 2. Wir machen für dieselbe von den in 4. angegebenen Bezeichnungen Gebrauch. Der Modul 2^t sei größer als die Potenz $2^{1+\omega_{n-1}}$; man hat dann immer $t \geq 2$, und wenn $\omega_{n-1} > 0$, auch $t \geq 3$.

Wir werden die Größen $f(2^t)$ finden, indem wir f als *Hauptrest* für den Modul 2^t annehmen, also für f den Typus zulassen:

$$f \equiv \{ \Phi_1 + 2^{\omega_{\mathfrak{I}_1}} [\Phi_2 + \dots + 2^{\omega_{\mathfrak{I}_{\lambda-1}}} (\Phi_{\lambda}) \dots] \} \pmod{2^t},$$

wo die einzelnen Φ_k Reste vorstellen, die in bezug auf 2 primitiv sind, und entweder dem Typus

$$(R_I) \quad \Phi_k \equiv \begin{pmatrix} \alpha_1^{(k)}, \\ \alpha_2^{(k)}, \\ \dots \\ \alpha_{\kappa_k}^{(k)} \end{pmatrix} \pmod{2^{t-\omega_{\mathfrak{I}_k-1}}}$$

oder, wenn κ_k gerade ist, auch dem Typus

$$(R_{II}) \quad \Phi_k \equiv \begin{pmatrix} 2\alpha_1^{(k)}, & A_1^{(k)}, \\ A_1^{(k)}, & 2\tilde{\alpha}_1^{(k)}, \\ & \dots \\ & \dots \\ & 2\alpha_{\frac{\kappa_k}{2}}^{(k)}, & A_{\frac{\kappa_k}{2}}^{(k)} \\ & A_{\frac{\kappa_k}{2}}^{(k)}, & 2\tilde{\alpha}_{\frac{\kappa_k}{2}}^{(k)} \end{pmatrix} \pmod{2^{t-v_{\mathfrak{g}_k-1}}}$$

angehören (*F. Q.*, p. 23 [[S. 25]]). Dabei bedeuten die $\alpha_h^{(k)}$, ebenso wie die $A_h^{(k)}$, lauter ungerade Größen.

Wir geben jedem Reste Φ_k vom Typus (R_I) eine Zahl $\tau_k = 1$ und jedem Reste Φ_k vom Typus (R_{II}) eine Zahl $\tau_k = 2$. Die $\kappa_k - 1$ Invarianten σ eines Restes Φ_k nennen wir $\varrho_1^{(k)}, \varrho_2^{(k)}, \dots, \varrho_{\kappa_k-1}^{(k)}$. Es gelten dann folgende Beziehungen (*F. Q.*, art. IV [[S. 28—29]]): wenn $\tau_k = 1$:

$$\varrho_h^{(k)} = 1;$$

wenn $\tau_k = 2$:

$$\varrho_{2h_0-1}^{(k)} = 2, \quad \varrho_{2h_0}^{(k)} = 1;$$

ferner:

$$\sigma_{\mathfrak{g}_{k-1+h}} = \varrho_h^{(k)}, \quad (h = 1, 2, \dots, \kappa_k - 1)$$

und:

$$\sigma_{\mathfrak{g}_k} = 1,$$

so daß die Zahlen τ_k durch die Invarianten σ_h bestimmt sind. In der Tat erhält man insbesondere:

$$\tau_k = \sigma_{\mathfrak{g}_{k-1}+1} = \sigma_{\mathfrak{g}_k-1}.$$

Ein Ausdruck von der Gestalt $\frac{1}{\tau}\Phi$ mag, je nachdem $\tau = 1$ oder $\tau = 2$ ist, kurz mit (I) oder mit (II) bezeichnet werden. Wir schicken zunächst einige Bemerkungen über die Kongruenzen

$$(I) \equiv m \pmod{4} \quad \text{und} \quad (II) \equiv m \pmod{2}$$

voraus.

I. Für einen Ausdruck

$$\Psi \equiv \alpha_1 \xi_1^2 + \alpha_2 \xi_2^2 + \dots + \alpha_{\kappa} \xi_{\kappa}^2 \pmod{4}$$

mögen $\Psi_0, \Psi_1, \Psi_2, \Psi_3$ die Anzahlen der Lösungen von

$$\Psi \equiv 0, 1, 2, 3 \pmod{4}$$

vorstellen. Diese 4 Anzahlen können leicht nach *F. Q.*, art. VIII [[S. 50 bis 58]], gefunden werden. Man setze nämlich

$$4\Psi_0 = \psi_0 + \psi_1 + \psi_2 + \psi_3,$$

$$4\Psi_1 = \psi_0 - i\psi_1 - \psi_2 + i\psi_3,$$

$$4\Psi_2 = \psi_0 - \psi_1 + \psi_2 - \psi_3,$$

$$4\Psi_3 = \psi_0 + i\psi_1 - \psi_2 - i\psi_3,$$

wo $i = \sqrt{-1}$, und bilde die Einheiten:

$$(6_I) \quad \varepsilon = (-1)^{\left[\frac{x}{2}\right]} \cdot (-1)^{\sum_{k=1}^x \frac{\alpha_k - 1}{2}},$$

$$\delta = (-1)^{\left[\frac{x}{4}\right]} \cdot (-1)^{\left[\frac{x}{2}\right] \left(\left[\frac{x}{2}\right] + \sum_{k=1}^x \frac{\alpha_k - 1}{2}\right)} \cdot (-1)^{\sum_{k'=1}^{1,x} \frac{\alpha_{k'} - 1}{2} \cdot \frac{\alpha_{k''} - 1}{2}}$$

Als dann geben die Formeln F , Q , p. 62 und p. 64 [[S. 55 und S. 57]]:

$$\psi_0 = 2^{2x}, \quad \psi_2 = 0$$

und

$$\psi_h = \left(\frac{1+i}{2}\varepsilon - \frac{1-i}{2}\varepsilon i\right) \varepsilon^{\frac{h-1}{2}} \cdot \delta \cdot (-i)^{x^2 \left(\frac{h-1}{2}\right)^2} \cdot \left(\frac{1+i}{\sqrt{2}}\right)^{x-2 \left[\frac{x}{2}\right]} \cdot 2^{\frac{3x}{2}}. \quad (h=1, 3)$$

Wir unterscheiden die folgenden Fälle:

1. $x \equiv 0 \pmod{2}$.

$$A) \quad \varepsilon = 1; \quad \psi_1 = \delta \cdot 2^{\frac{3x}{2}}, \quad \psi_3 = \delta \cdot 2^{\frac{3x}{2}}.$$

$$4\Psi_0 = 2^{2x} + \delta \cdot 2^{\frac{3x}{2}+1},$$

$$4\Psi_2 = 2^{2x} - \delta \cdot 2^{\frac{3x}{2}+1},$$

$$4\Psi_1 = 4\Psi_3 = 2^{2x}.$$

$$B) \quad \varepsilon = -1; \quad \psi_1 = -i\delta \cdot 2^{\frac{3x}{2}}, \quad \psi_3 = i\delta \cdot 2^{\frac{3x}{2}}.$$

$$4\Psi_0 = 4\Psi_2 = 2^{2x},$$

$$4\Psi_1 = 2^{2x} - \delta \cdot 2^{\frac{3x}{2}+1},$$

$$4\Psi_3 = 2^{2x} + \delta \cdot 2^{\frac{3x}{2}+1}.$$

2. $x \equiv 1 \pmod{2}$.

$$A) \quad \varepsilon = 1; \quad \psi_1 = \delta \cdot \frac{1+i}{\sqrt{2}} \cdot 2^{\frac{3x}{2}}, \quad \psi_3 = \delta \cdot \frac{1-i}{\sqrt{2}} \cdot 2^{\frac{3x}{2}}.$$

$$4\Psi_0 = 4\Psi_1 = 2^{2x} + \delta \cdot 2^{\frac{3x+1}{2}},$$

$$4\Psi_2 = 4\Psi_3 = 2^{2x} - \delta \cdot 2^{\frac{3x+1}{2}}$$

$$B) \quad \varepsilon = -1; \quad \psi_1 = \delta \cdot \frac{1-i}{\sqrt{2}} \cdot 2^{\frac{3x}{2}}, \quad \psi_3 = \delta \cdot \frac{1+i}{\sqrt{2}} \cdot 2^{\frac{3x}{2}}.$$

$$4\Psi_0 = 4\Psi_3 = 2^{2x} + \delta \cdot 2^{\frac{3x+1}{2}},$$

$$4\Psi_2 = 4\Psi_1 = 2^{2x} - \delta \cdot 2^{\frac{3x+1}{2}}.$$

In betreff der Einheiten ε und δ erwähnen wir noch einige Punkte.

1. Der Rest $l \equiv \alpha_1 + \alpha_2 + \dots + \alpha_x \pmod{4}$ ist durch die Einheit ε bestimmt.

Da nämlich immer $\alpha_k \equiv 1 \pmod{2}$ ist, so kommt zunächst

$$l \equiv x \pmod{2}, \quad (-1)^l = (-1)^x.$$

Sind ferner von den Zahlen $\alpha_1, \alpha_2, \dots, \alpha_x$ im ganzen x_0 kongruent $-1 \pmod{4}$ und $x - x_0$ kongruent $1 \pmod{4}$, so folgt einerseits:

$$\varepsilon = (-1)^{\left[\frac{x}{2}\right] - x_0},$$

andererseits:

$$l \equiv (x - x_0) - x_0 \equiv x - 2x_0 \pmod{4},$$

mithin:

$$(-1)^{\left[\frac{l}{2}\right]} = (-1)^{\left[\frac{x}{2}\right] - x_0} = \varepsilon.$$

Im speziellen findet man, wenn $x \equiv 1 \pmod{2}$:

$$l \equiv \varepsilon \pmod{4}.$$

2. Ersetzt man Ψ durch $-\Psi$, also $\alpha_k \pmod{4}$ durch $-\alpha_k \pmod{4}$, so mögen an die Stelle von ε und δ die Einheiten ε^- und δ^- treten. Man erhält unmittelbar:

$$\varepsilon^- = \varepsilon \cdot (-1)^x, \quad \delta^- = \delta \cdot \varepsilon^{x-1},$$

also

$$1) \ x \equiv 0 \pmod{2}: \quad \varepsilon^- = \varepsilon, \quad \delta^- = \delta \cdot \varepsilon;$$

$$2) \ x \equiv 1 \pmod{2}: \quad \varepsilon^- = -\varepsilon, \quad \delta^- = \delta.$$

Dasselbe Resultat erschließt man auch leicht aus der aufgestellten Tabelle, indem man die Beziehungen $\Psi_{-h} = (-\Psi)_h$ beachtet.

3. Wir wollen

$$\Psi^1 \equiv \alpha_2 \xi_2^2 + \dots + \alpha_x \xi_x^2 \pmod{4}$$

setzen und die Einheiten ε und δ , welche zu Ψ^1 gehören, mit ε^1 und δ^1 bezeichnen.

Mit Hilfe der Relationen

$$\left[\frac{x}{2}\right] - \left[\frac{x-1}{2}\right] \equiv x-1, \quad \left[\frac{x}{4}\right] - \left[\frac{x-1}{4}\right] \equiv (x-1) \left[\frac{x-1}{2}\right] \pmod{2}$$

findet man:

$$\varepsilon = (-1)^{x-1} \cdot (-1)^{\frac{\alpha-1}{2}} \cdot \varepsilon^1, \quad \delta = (-1)^{x \cdot \frac{\alpha-1}{2}} \cdot \varepsilon^{(x-1) + \frac{\alpha-1}{2}} \cdot \delta^1, \quad (\alpha = \alpha_1)$$

also

1. $x \equiv 0 \pmod{2}$;

$$A) \ \varepsilon = 1: \quad \delta = \delta^1, \quad \alpha \equiv -\varepsilon^1 \pmod{4};$$

$$B) \ \varepsilon = -1: \quad \delta = -(-1)^{\frac{\alpha-1}{2}} \cdot \delta^1, \quad \alpha \equiv \varepsilon^1 \pmod{4}.$$

2. $x \equiv 1 \pmod{2}$;

$$A) \ \alpha \equiv -\varepsilon \pmod{4}: \quad \varepsilon^1 = -1, \quad \delta = -\varepsilon \cdot \delta^1;$$

$$B) \ \alpha \equiv \varepsilon \pmod{4}: \quad \varepsilon^1 = 1, \quad \delta = \delta^1.$$

II. Jetzt liege ein Ausdruck vor:

$$X \equiv (\alpha_1 \xi_1^2 + A_1 \xi_1 \tilde{\xi}_1 + \tilde{\alpha}_1 \tilde{\xi}_1^2) + \cdots + \left(\alpha_{\frac{x}{2}} \xi_{\frac{x}{2}}^2 + A_{\frac{x}{2}} \xi_{\frac{x}{2}} \tilde{\xi}_{\frac{x}{2}} + \tilde{\alpha}_{\frac{x}{2}} \tilde{\xi}_{\frac{x}{2}}^2 \right) \pmod{2},$$

und es bedeute X_0 und X_1 die Anzahl der Lösungen von

$$X \equiv 0 \quad \text{und} \quad X \equiv 1 \pmod{2}.$$

Setzt man

$$2X_0 = \chi_0 + \chi_1, \quad 2X_1 = \chi_0 - \chi_1,$$

ferner:

$$(6_{II}) \quad \theta = \left(\frac{2}{4\alpha_1 \tilde{\alpha}_1 - A_1^2} \right) \cdots \left(\frac{2}{4\alpha_{\frac{x}{2}} \tilde{\alpha}_{\frac{x}{2}} - A_{\frac{x}{2}}^2} \right),$$

so kommt nach *F. Q.*, p. 65 [[S. 58]]:

$$\chi_0 = 2^*, \quad \chi_1 = \theta \cdot 2^{\frac{x}{2}},$$

also:

$$2X_0 = 2^* + \theta \cdot 2^{\frac{x}{2}}, \quad 2X_1 = 2^* - \theta \cdot 2^{\frac{x}{2}}.$$

Nach diesen Vorbereitungen wollen wir versuchen, eine Formel aufzustellen, mit deren Hilfe die Größen $f(2')$ für Reste von $n(>1)$ Variablen auf entsprechende Größen für Reste von weniger als n Variablen zurückgeführt werden können. Für Reste f von einer Variablen hat man einfach $f(2') = 1$.

Da wir f als Hauptrest für den Modul $2'$ voraussetzen, so gilt jedenfalls eine Kongruenz

$$f \equiv \Phi_1 + 2^{\omega_{x_1}} f(x_1) \pmod{2'}, \quad (\omega_{x_1} \geq 1)$$

wo Φ_1 einen Rest vom Typus (R_I) oder (R_{II}) bedeutet. Wir unterscheiden zwei Fälle, je nachdem die Invariante $\tau_1 = \sigma_1$ den Wert 1 oder 2 erhält.

$$(\sigma_1 = 1).$$

Ist zunächst $\sigma_1 = 1$, so läßt f sich zugleich in der Form schreiben:

$$f \equiv \alpha \xi^2 + 2^{\omega_1} f^{(1)} \pmod{2'},$$

wo α ungerade ist und $f^{(1)}$ einen in bezug auf 2 primitiven Rest vorstellt, welcher im Falle $\omega_1 = 0$ ($x_1 > 1$) eine erste Invariante σ gleich 1 ergibt. Überhaupt folgen die Invarianten $\sigma_h^{(1)}$ und $2^{\omega_h^{(1)}}$ von $f^{(1)}$ aus den Gleichungen:

$$\sigma_{h-1}^{(1)} = \sigma_h, \quad \omega_{h-1}^{(1)} = \omega_h. \quad (h=2, 3, \dots, n-1)$$

Die Anzahl $f(2')$ der inkongruenten Substitutionen T von einer Determinante $\equiv 1 \pmod{2'}$, welche den Rest f in sich selbst überführen, drückt sich in ähnlicher Weise aus wie oben die Zahl $f(p')$. In der Tat, die erste Vertikalreihe einer Substitution T muß immer von n Zahlen $\xi_i \pmod{2'}$ gebildet sein, welche

$$(3) \quad f(\xi_i) \equiv \alpha \pmod{2'}$$

liefern. Dazu tritt in den Fällen, wo $\kappa_1 > 1$ ist, noch die Bedingung, daß nicht zu gleicher Zeit die sämtlichen Kongruenzen

$$(7) \quad \xi_1 \equiv 1, \xi_2 \equiv 1, \dots, \xi_{\kappa_1} \equiv 1 \pmod{2}$$

bestehen dürfen (*F. Q.*, p. 172 und 128 [[S. 139 und 106]]). Wir bezeichnen mit $A \cdot 2^{(n-1)t}$, je nachdem $\kappa_1 = 1$ oder > 1 ist, entweder die Anzahl aller möglichen Lösungen ξ_i von (3) oder nur die Anzahl aller derjenigen Lösungen ξ_i , welche nicht zugleich die Kongruenzen (7) erfüllen.

Die Betrachtungen in *F. Q.*, p. 172 [[S. 139]], zeigen, daß wirklich jedem dieser Systeme ξ_i Substitutionen T entsprechen, welche den Rest f in sich selbst transformieren. Ebenso wie in 5. schließt man dann, daß jedes solche System ξ_i zu genau soviel verschiedenen T Veranlassung gibt, als verschiedene Substitutionen von einer Determinante $\equiv 1 \pmod{2^t}$ existieren, welche auf den Rest $2^{\omega_1} f^{(1)}$ ohne Wirkung sind. So gewinnt man die Beziehung:

$$f(2^t) = 2^{(n-1)t + \omega_1[(n-1)^2 - 1]} \cdot A \cdot f^{(1)}(2^{t-\omega_1}).$$

In betreff der Größe A unterscheiden wir die Fälle $\kappa_1 = 1$ und $\kappa_1 > 1$.

1. Ist zunächst $\kappa_1 = 1$, so gibt unsere Annahme über t (wenn $n > 1$) jedenfalls $t \geq 3$. Nach dem Satze 4. (C) muß daher $A \cdot 2^{8(n-1)}$ die Anzahl der Lösungen von $f(\xi_i) \equiv \alpha \pmod{8}$ ausdrücken.

Indem man in f alle Glieder fortläßt, welche durch 8 teilbar sind, erlangt f entweder den Typus:

$$(8) \quad f \equiv \alpha \xi^2 \pmod{8}.$$

In diesem Falle hat man $f \equiv \alpha \pmod{8}$, sobald $\xi \equiv 1 \pmod{2}$ ist. Die Anzahl der Lösungen von $f \equiv \alpha \pmod{8}$ beträgt demnach $4 \cdot 2^{8(n-1)}$ und man findet:

$$A = 4.$$

Oder f gehört einem der beiden Typen an:

$$(9) \quad \begin{aligned} f &\equiv \alpha \xi^2 + 4(I) \\ f &\equiv \alpha \xi^2 + 4(II) + 4(I) \end{aligned} \pmod{8}.$$

In dem Ausdrucke 4(I) erscheint mindestens ein Glied $4\alpha \mathfrak{x}^2 \equiv 4\mathfrak{x} \pmod{8}$. Durch Änderung des Restes von $\mathfrak{x} \pmod{2}$ können wir aus jeder Lösung von $f \equiv \alpha \pmod{8}$ eine solche von $f \equiv \alpha + 4 \pmod{8}$ herleiten, und umgekehrt. Diese zwei Kongruenzen besitzen also gleich viel, nämlich $2 \cdot 2^{8(n-1)}$ Lösungen, und man erhält:

$$A = 2.$$

Oder f gehört dem Typus an:

$$(10) \quad f \equiv \alpha \xi^2 + 4(II) \pmod{8}.$$

In diesem Falle hat $f \equiv \alpha \pmod{8}$ stets $\xi \equiv 1 \pmod{2}$ und (II) $\equiv 0 \pmod{2}$ zur Folge. Bildet man für den Rest (II) von $\kappa_2 = \kappa$ Variablen,

nach der Formel (6_{II}), eine Einheit θ , so ergibt sich die Anzahl der Lösungen von $(II) \equiv 0 \pmod{2}$ gleich $2^{\kappa-1} \left(1 + \theta \cdot 2^{\frac{-\kappa}{2}}\right)$, und man bekommt

$$A = 2 \left(1 + \frac{\theta}{2^{\frac{\kappa}{2}}}\right).$$

Oder f gehört dem Typus an:

$$(11) \quad f \equiv \alpha \xi^2 + 2(I) \pmod{8}.$$

Dann ist $f \equiv \alpha \pmod{8}$ identisch mit $\xi \equiv 1 \pmod{2}$ und $(I) \equiv 0 \pmod{4}$. Gehören zu dem Reste (I) von $\kappa_2 = \kappa$ Variablen, gemäß der Formel (6_I), zwei Einheiten ε und δ , so liefert die obige Tabelle:

$$1) \quad \kappa \equiv 0 \pmod{2},$$

$$\varepsilon = 1: \quad A = \left(1 + \frac{\delta}{2^{\frac{\kappa}{2}-1}}\right);$$

$$\varepsilon = -1: \quad A = 1.$$

$$2) \quad \kappa \equiv 1 \pmod{2},$$

$$A = \left(1 + \frac{\delta}{2^{\frac{\kappa-1}{2}}}\right).$$

Oder f gehört endlich dem Typus an:

$$(12) \quad f \equiv \alpha \xi^2 + 2(I) + 4(I)_1 \pmod{8}.$$

In diesem Falle enthält $2(I)$ mindestens ein Glied $2\alpha \chi^2 \equiv 2\chi \pmod{4}$, mit dessen Hilfe die Lösungen von $f \equiv 1$ und $f \equiv 3 \pmod{4}$ einander eindeutig zugeordnet werden können; und ebenso erscheint in $4(I)_1$ mindestens ein Glied $4\alpha \chi^2 \equiv 4\chi \pmod{8}$, infolge dessen die Kongruenzen $f \equiv \alpha$ und $f \equiv \alpha + 4 \pmod{8}$ gleich viel Lösungen zulassen. So findet man leicht:

$$A = 1.$$

2. Jetzt sei $\kappa_1 > 1$. Wir teilen für einen Moment die Systeme ξ_i , welche $f(\xi_i) \equiv 1 \pmod{2}$ ergeben, in Systeme erster oder zweiter Art ein, je nachdem sie den Bedingungen (7) entgegen sind oder mit denselben harmonieren. Irgendein System $\xi_i \pmod{4}$ erster Art möge die Kongruenz

$$(13) \quad f(\xi_i) \equiv \alpha \pmod{4}$$

erfüllen. Da ein solches System zugleich einen bestimmten Wert des Restes $f(\xi_i) \pmod{8}$ ergibt, so wird es nur einer der beiden Kongruenzen $f(\xi_i) \equiv \alpha \pmod{8}$ und $f(\xi_i) \equiv \alpha + 4 \pmod{8}$ Genüge leisten. Ist nun von den Zahlen $\xi_1, \xi_2, \dots, \xi_{\kappa_1}$ etwa ξ_k die erste, welche nicht ungerade ausfällt, und ändern wir den Rest $\xi_k \pmod{4}$ in $\xi_k + 2 \pmod{4}$, so geht aus dem Systeme $\xi_i \pmod{4}$ ein anderes System erster Art hervor, welches offenbar der anderen von diesen beiden Kongruenzen genügen muß. So erhellt, daß diese zwei Kongruenzen gleich viel Lösungen erster Art zulassen.

Beachtet man jetzt die Definition der Größe A und den Satz 4. (C), so folgt, daß in allen Fällen die Anzahl der Lösungen erster Art von (13) durch $A \cdot 2^{2(\kappa-1)}$ ausgedrückt ist. Man gelangt zu dieser Anzahl, indem man die Anzahl aller möglichen Lösungen dieser Kongruenz um die Anzahl ihrer Lösungen zweiter Art vermindert.

Wir unterscheiden die Fälle $\kappa_1 \equiv 0 \pmod{2}$ und $\kappa_1 \equiv 1 \pmod{2}$, schreiben aber der Einfachheit halber κ für κ_1 .

1°. Zunächst sei $\kappa \equiv 0 \pmod{2}$. In diesem Falle treten überhaupt keine Lösungen zweiter Art auf, da die Kongruenzen (7) stets $f(\xi_i) \equiv \kappa \pmod{2}$ nach sich ziehen.

Entweder ist nun f von dem Typus:

$$(14) \quad f \equiv (I) \pmod{4}.$$

Lassen wir dieselben Bezeichnungen wie oben für Ψ gelten, so liefert die aufgestellte Tabelle:

$$\begin{aligned} \varepsilon = 1: \quad A = 1; \quad & [\delta = \delta^1, (-1)^{\frac{\alpha-1}{2}} = -\varepsilon^1] \\ \varepsilon = -1: \quad A = \left(1 + \frac{\delta^1}{2^{\frac{\kappa-1}{2}}}\right); \quad & [\delta = -(-1)^{\frac{\alpha-1}{2}} \cdot \delta^1, (-1)^{\frac{\alpha-1}{2}} = \varepsilon^1] \end{aligned}$$

Oder f ist von dem Typus:

$$(15) \quad f \equiv (I) + 2(I)_1 \pmod{4}.$$

Alsdann erscheint in $2(I)_1$ ein Glied $2\alpha\chi^2 \equiv 2\chi \pmod{4}$, welches bewirkt, daß $f \equiv 1$ und $f \equiv 3 \pmod{4}$ gleich viel Lösungen zulassen. Demnach kommt einfach:

$$A = 1.$$

2°. Endlich sei $\kappa \equiv 1 \pmod{2}$. Wir unterscheiden dieselben zwei Fälle.

Entweder ist f von dem Typus:

$$(14) \quad f \equiv (I) \pmod{4}.$$

Identifizieren wir den Rest (I) mit dem obigen Ausdrucke Ψ , so entsteht für Systeme ξ_i zweiter Art:

$$f(\xi_i) \equiv \alpha_1 + \alpha_2 + \cdots + \alpha_\kappa \equiv \varepsilon \pmod{4}.$$

Diese Systeme gehen also nur den Fall an, wo $\alpha \equiv \varepsilon \pmod{4}$ ist. Mithin kommt:

$$\alpha \equiv -\varepsilon \pmod{4}: \quad A = \left(1 - \frac{\delta}{2^{\frac{\kappa-1}{2}}}\right); \quad [\varepsilon^1 = -1, \delta = -\varepsilon \cdot \delta^1]$$

$$\alpha \equiv \varepsilon \pmod{4}: \quad A = \left(1 + \frac{\delta}{2^{\frac{\kappa-1}{2}}} - \frac{1}{2^{\frac{\kappa-2}{2}}}\right) = \left(1 - \frac{\delta}{2^{\frac{\kappa-1}{2}}}\right) \left(1 + \frac{\delta^1}{2^{\frac{\kappa-3}{2}}}\right).$$

$$[\varepsilon^1 = 1, \delta = \delta^1]$$

Oder f ist vom Typus:

$$(15) \quad f \equiv (I) + 2(I)_1 \pmod{4}.$$

Alsdann bewirkt ein jedes Glied $2\alpha x^2 \equiv 2x \pmod{4}$ aus $2(I)_1$, daß $f \equiv 1$ und $f \equiv 3 \pmod{4}$ sowohl gleich viel Lösungen erster, wie gleichviel Lösungen zweiter Art besitzen, und man findet:

$$A = \left(1 - \frac{1}{2^{x-1}}\right) \cdot \\ (\sigma_1 = 2).$$

Wenn $\sigma_1 = 2$ ist, so läßt f sich zugleich in der Form schreiben:

$$f \equiv 2(\alpha \xi^2 + A \xi \tilde{\xi} + \tilde{\alpha} \tilde{\xi}^2) + 2^{\omega_2} f^{(2)} \pmod{2^t},$$

wo α und A ungerade sind, und $f^{(2)}$ einen in bezug auf 2 primitiven Rest bedeutet. Die Invarianten $\sigma_h^{(2)}$ und $2^{\omega_h^{(2)}}$ von $f^{(2)}$ bestimmen sich aus den Gleichungen:

$$\sigma_{h-2}^{(2)} = \sigma_h, \quad \omega_{h-2}^{(2)} = \omega_h. \quad (h = 3, 4, \dots, n-1)$$

Wir betrachten die $f^{(2^t)}$ Substitutionen T von einer Determinante $\equiv 1 \pmod{2^t}$, welche den Rest f in sich selbst verwandeln. In jeder dieser Substitutionen muß die erste Vertikalreihe von n Zahlen $\xi_i \pmod{2^t}$ gebildet werden, welche

$$(3) \quad f(\xi_i) \equiv 2\alpha \pmod{2^t}$$

ergeben. Ferner dürfen diese Zahlen nicht zugleich alle Bedingungen

$$(7) \quad \xi_1 \equiv 0, \xi_2 \equiv 0, \dots, \xi_{n_1} \equiv 0 \pmod{2}$$

erfüllen (*F. Q.*, p. 175 und 129 [[S. 141 und 107]]). Wir bezeichnen mit $2A \cdot 2^{(n-1)t}$ die Anzahl aller Systeme $\xi_i \pmod{2^t}$, welche der Kongruenz (3), aber nicht sämtlichen Kongruenzen (7) genügen.

Jedes dieser Systeme ξ_i tritt wirklich in mindestens einer von den Substitutionen T als erste Vertikalreihe auf (*F. Q.*, pp. 175—177 [[S. 141—142]]), also etwa in einem T_0 . Es fragt sich, für wieviel verschiedene T ein solches System ξ_i die erste Vertikalreihe bildet; offenbar für alle diejenigen T , welche zusammengesetzt sind aus T_0 und aus einer Substitution T' mit der ersten Vertikalreihe $1, 0, \dots, 0$. Man hat also zu ermitteln, wie viele Substitutionen vom Typus

$$T \equiv \begin{vmatrix} 1, & U & , & \dots & \dots \\ 0, & \eta_0 & , & \dots & \dots \\ 0, & \eta_1 & , & \dots & \dots \\ . & . & . & . & . \\ 0, & \eta_{n-2}, & \dots & \dots & \dots \end{vmatrix} \equiv 1 \pmod{2^t}$$

den Rest f in sich selbst transformieren.

Setzen wir

$$f^{(1)} \equiv (4\alpha\tilde{\alpha} - A^2)\eta_0^2 + 2^{\omega_2+1} \cdot \alpha \cdot f^{(2)}(\eta_1, \dots, \eta_{n-2}) \pmod{2^{t+1}},$$

so entsteht zunächst die Bedingung

$$f^{(1)}(\eta_0, \eta_1, \dots, \eta_{n-2}) \equiv (4\alpha\tilde{\alpha} - A^2) \pmod{2^{t+1}}.$$

Dieser Kongruenz mögen $\frac{1}{2} A^{(1)} \cdot 2^{(n-2)t}$ verschiedene Systeme $\eta_i \pmod{2^t}$ Genüge leisten. Durch Überlegungen, wie sie in *F. Q.*, pp. 176—177 [[S. 141—142]], angestellt sind, überzeugt man sich, daß wirklich jedem dieser Systeme η_i Substitutionen T zukommen, welche den Rest f in sich selbst transformieren; es fragt sich wieviele.

Die Gesamtheit aller solcher T wird hervorgehen, indem man eine beliebige unter ihnen, etwa T_0 , mit allen Substitutionen

$$\mathfrak{Z} \equiv \begin{vmatrix} 1, & U, & V_1, & \dots, & V_{n-2} \\ 0, & 1, & \tilde{V}_1, & \dots, & \tilde{V}_{n-2} \\ 0, & 0, & t_1^1, & \dots, & t_1^{n-2} \\ \cdot & \cdot & \cdot & \cdot & \cdot \\ 0, & 0, & t_{n-2}^1, & \dots, & t_{n-2}^{n-2} \end{vmatrix} \equiv 1 \pmod{2^t}$$

zusammensetzt, welche $f \pmod{2^t}$ in sich selbst überführen. Für $\mathfrak{Z}$ findet man $2U \equiv 0$, $V_h \equiv 0$, $\tilde{V}_h \equiv 0 \pmod{2^t}$; und man erkennt, daß die Anzahl der verschiedenen $\mathfrak{Z}$ durch $2F(2^t)$ ausgedrückt ist, falls F den Rest $2\omega_2 f^{(2)}$ bedeutet. So ergibt sich die Beziehung

$$f(2^t) = 2^{(n-1)t + (n-2)t + 1 + \omega_2[(n-2)^2 - 1]} \cdot A \cdot A^{(1)} \cdot f^{(2)}(2^{t-\omega_2}).$$

Nun hat man zugleich

$$f^{(1)}(2^{t+1}) = 2^{(n-2)(t+1) + (\omega_2+1)[(n-2)^2 - 1]} \cdot A^{(1)} \cdot f^{(2)}(2^{t-\omega_2});$$

also kommt endlich:

$$f(2^t) = 2^{(n-1)t - n(n-3)} \cdot A \cdot f^{(1)}(2^{t+1}).$$

Um die Größe A zu finden, beachte man, daß das System der κ_1 Kongruenzen (7) sich als identisch mit dem Systeme $\frac{1}{2} \frac{\partial f}{\partial \xi_i} \equiv 0 \pmod{2}$ ($i = 1, 2, \dots, n$) erweist. Nach 4. (D) muß infolgedessen $A \cdot 2^{n-1}$ die Anzahl aller Lösungen von

$$(13) \quad \frac{1}{2} f(\xi_i) \equiv 1 \pmod{2}$$

vorstellen, bei welchen die n Zahlen $\frac{1}{2} \frac{\partial f}{\partial \xi_i}$ nicht sämtlich gerade sind. Wir wollen für κ_1 einfach κ setzen.

Entweder ist $\frac{1}{2} f$ vom Typus:

$$(14) \quad \frac{1}{2} f \equiv (\text{II}) \pmod{2}.$$

In diesem Falle liefern die Kongruenzen (7) stets $f \equiv 0 \pmod{4}$; sie sind also nicht mit (13) verträglich. Gehört zu (II) wie oben eine Einheit θ , so kommt demnach:

$$A = \left(1 - \frac{\theta}{2^{\frac{x}{2}}}\right).$$

Oder $\frac{1}{2}f$ ist vom Typus:

$$[15] \quad \frac{1}{2}f \equiv (\text{II}) + (\text{I}) \pmod{2}.$$

Dann erscheint in (I) ein Glied $\alpha x^2 \equiv x \pmod{2}$, welches zur Folge hat, daß $\frac{1}{2}f \equiv 0$ und $\frac{1}{2}f \equiv 1 \pmod{2}$ gleich viel Lösungen zulassen, man betrachte diese Kongruenzen für sich oder zusammen mit den Kongruenzen (7). Man erhält also:

$$A = \left(1 - \frac{1}{2^x}\right).$$

Die Größe $A^{(1)}$ bestimmt sich mit Hilfe der Formeln aus $(\sigma_1 = 1)$ Im Falle [14] findet man, wenn $x = 2$: $A^{(1)} = 4$, und wenn $x > 2$:

$$A^{(1)} = 2 \left(1 + \frac{\theta^1}{2^{\frac{x}{2}-1}}\right), \quad \text{wobei} \quad \theta = \left(\frac{2}{4\alpha\bar{\alpha} - A^2}\right) \cdot \theta^1;$$

im Falle [15] wird immer $A^{(1)} = 2$.

Wir setzen jetzt allgemein:

$$f(2^t) = 2^{\frac{n(n-1)}{2}t + \sum_{h=1}^{n-1} \omega_h \left(\frac{(n-h)(n-h+1)}{2} - 1\right)} \cdot \prod_{h=1}^{n-1} \sigma_h \cdot f\{2\}. \quad \left(t > 1 + \sum_{h=1}^{n-1} \omega_h\right)$$

Unsere Rekursionsformeln gehen dann in

$$f\{2\} = A \cdot f^{(1)}\{2\}$$

über, und auf Grund der gefundenen Werte von A können wir den vollständigen Ausdruck der Größe $f\{2\}$ hinschreiben.

Es bedeute, wie bisher, f einen aus λ Gliedern bestehenden Hauptrest, wo λ gleich ist der um 1 vermehrten Anzahl aller durch 2 teilbaren Größen 2^{ω_h} ($h = 1, \dots, n-1$). Ferner bezeichne $\mu - 1$ die Anzahl aller Größen aus der Reihe $\sigma_{h-1} 2^{\omega_h} \sigma_{h+1}$ ($h = 1, \dots, n-1$), welche den Faktor 4 enthalten, und $\nu - 1$ die Anzahl aller Größen dieser Reihe welche durch 8 aufgehen.

Die Größe $f\{2\}$ ist dann gleich einem Produkte aus der Potenz $\frac{2^{(2\mu-1)+(\nu-1)}}{\prod \sigma_h}$ und aus λ Faktoren $\mathfrak{A}_k$, welche den λ einzelnen Resten Φ_k entsprechen und folgendermaßen bestimmt werden:

(I). So oft Φ_k ein Rest vom Typus (R_I) ist, nehme man:

$$\mathfrak{A}_k = \left(1 - \frac{1}{2^2}\right) \left(1 - \frac{1}{2^4}\right) \cdots \left(1 - \frac{1}{2^{\left[\frac{x_k-1}{2}\right]}}\right) \cdot a_k,$$

und setze $\alpha_k = 1$, falls die Zahlen $\tau_{k-1} \cdot 2^{\omega_{\mathfrak{g}_k-1}}$ und $2^{\omega_{\mathfrak{g}_k}} \cdot \tau_{k+1}$ nicht beide durch 4 teilbar sind; andernfalls aber bilde man für Φ_k , gemäß den Formeln (6_I), zwei Einheiten ε_k und δ_k , und setze:

1) wenn $\kappa_k \equiv 0 \pmod{2}$,

je nachdem $\varepsilon_k = 1$ oder $= -1$ ist,

$$\alpha_k = \left(1 + \delta_k \cdot 2^{-\frac{\kappa_k}{2} + 1}\right)^{-1} \quad \text{oder} \quad = 1;$$

2) wenn $\kappa_k \equiv 1 \pmod{2}$,

$$\alpha_k = \left(1 + \delta_k \cdot 2^{-\frac{\kappa_k-1}{2}}\right)^{-1}.$$

(II). So oft Φ_k ein Rest vom Typus (R_{II}) ist, nehme man:

$$\mathfrak{A}_k = \left(1 - \frac{1}{2^2}\right) \left(1 - \frac{1}{2^4}\right) \cdots \left(1 - \frac{1}{2^{\kappa_k}}\right) \cdot \alpha_k,$$

und setze $\alpha_k = 1$, falls die Zahlen $\tau_{k-1} \cdot 2^{\omega_{\mathfrak{g}_k-1}}$ und $2^{\omega_{\mathfrak{g}_k}} \cdot \tau_{k+1}$ nicht beide durch 4 teilbar sind; andernfalls aber bilde man für Φ_k , gemäß der Formel (6_{II}), eine Einheit θ_k , und setze:

$$\alpha_k = \left(1 + \theta_k \cdot 2^{-\frac{\kappa_k}{2}}\right)^{-1}.$$

Die Richtigkeit dieser Ausdrücke ergibt sich mit Hilfe eines Schlusses von $n-1$ auf n .

Wir erwähnen noch folgende Relation. Bedeutet $M-1$ die Anzahl aller durch 4 teilbaren Größen $\tau_k \cdot 2^{\omega_{\mathfrak{g}_k}} \cdot \tau_{k+1}$, so hat man

$$2^{\mu-1} = 2^{M-1} \cdot \frac{\sigma_1 \sigma_2 \cdots \sigma_{n-1}}{\tau_1 \tau_2 \cdots \tau_\lambda}.$$

Die Kongruenzen $f(\xi_i) \equiv m \pmod{2^i}$ sind sehr eingehend von Herrn C. Jordan in der Abhandlung *Sur la forme canonique des congruences du second degré et le nombre de leurs solutions**) untersucht worden. Die über diesen Gegenstand hier angestellten Betrachtungen sind indes wesentlich anderer Art. Die am Anfange gegebenen Werte der Zahlen Ψ_λ und X_λ wird man auch aus Art. 8 des *Mémoire sur la représentation des nombres par des sommes de cinq carrés***) von H. Smith ableiten können.

7. Die Größen $f\{q\}$ in ihrer Abhängigkeit von den Charakteren C .

Die Einheiten, welche in den Größen $f\{q\}$ auftreten, sind offenbar Invarianten der Form f . Sie müssen sich daher mit Hilfe der besonderen Charaktere C ausdrücken lassen, die wir in 2. aufgezählt haben.

Um dieses darzutun, setzen wir f , wie in 2., als charakteristische Form

*) Journal de Liouville, Deuxième Série, T. XVII, 1872, pp. 368—402.

**) Mémoires présentés à l'Académie des Sciences de Paris, T. XXIX, No. 1. (Collected Papers, vol. II, p. 623.)

ihres Genus voraus. Die aus den ersten h Reihen von f gebildeten symmetrischen Minoren mögen also Werte $\sigma_h d_{h-1} \varphi_h$ von solcher Art liefern, daß ein jedes φ_h relativ prim zu $2o_1 o_2 \dots o_{n-1}$ und zu $\varphi_{h-1} \cdot \varphi_{h+1}$ ausfällt. Ferner sei f primitiv.

Bedeutet q zunächst irgendeine ungerade Primzahl, die nicht in Δ aufgeht, so hängt $f\{q\}$, außer von der Zahl n , nur in dem Falle eines geraden n , noch von einer Einheit θ ab. Man findet dieselbe gleich

$$\left(\frac{(-1)^{\frac{n}{2}} \Delta}{q} \right) = \left(\frac{(-1)^{\frac{n}{2}-1} o_1 o_2 \dots o_{n-2} o_{n-1}}{q} \right).$$

Ist weiter $q = p$ irgendeine ungerade Primzahl aus Δ , so sind, laut Voraussetzung, sämtliche Zahlen φ_h zu p prim. Wir besitzen also in f eine *Grundform* für den Modul p (*F. Q.*, pp. 35–36 [[S. 34–35]]). Die Klasse f liefert infolgedessen für einen jeden Modul p^t unter anderen Resten auch folgenden Hauptrest:

$$\varphi \equiv \begin{pmatrix} \sigma_1 \varphi_1 \\ \sigma_0 \varphi_0, \\ \\ o_1 \cdot \sigma_2 \varphi_2 \\ \sigma_1 \varphi_1, \\ \\ \cdot \quad \cdot \quad \cdot \\ \cdot \quad \cdot \quad \cdot \\ \\ o_1 o_2 \dots o_{n-1} \cdot \frac{\sigma_n \varphi_n}{\sigma_{n-1} \varphi_{n-1}} \end{pmatrix} \pmod{p^t}.$$

In dem Ausdrucke von $f\{p\} = \varphi\{p\}$ gehört zu jeder von den λ Zahlen x_k , welche gerade ausfällt, eine Einheit θ . Setzt man $\vartheta_{k-1} = r$, $\vartheta_k = s$, und ist also $s - r = x_k \equiv 0 \pmod{2}$, so nimmt eine solche Einheit den Wert an:

$$\left(\frac{(-1)^{\frac{s-r}{2}} o_{r+1} o_{r+3} \dots o_{s-3} o_{s-1} \cdot \sigma_r \varphi_r \sigma_s \varphi_s}{p} \right).$$

Endlich sei $q = 2$. Nach Voraussetzung sind alle Zahlen φ_h ungerade, und f stellt eine *Grundform* für den Modul 2 vor. Man gewinnt daher, nach *F. Q.*, pp. 34–36 [[S. 33–35]] einen Hauptrest

$$\varphi \equiv \{ \Phi_1 + 2^{\omega \vartheta_1} [\Phi_2 + \dots + 2^{\omega \vartheta_{\lambda-1}} (\Phi_\lambda) \cdot \cdot] \} \pmod{2^t}$$

der Klasse f , indem man die Einzelreste Φ_k in folgender Weise auswählt. Zur Abkürzung sei $\vartheta_{k-1} = r$, $\vartheta_k = s$; dann nehme man, wenn $\tau_k = 1$:

$$\left. \begin{aligned} \frac{2^{\vartheta_r} \cdot \Phi_k}{o_1 o_2 \dots o_r} &\equiv \psi_1 + o_{r+1} \psi_2 + \dots + o_{r+1} o_{r+2} \dots o_{s-1} \psi_{s-r} \\ \psi_i &\equiv - \frac{\varphi_{r+i}}{\varphi_{r+i-1}} \xi_i^2 \end{aligned} \right\} \pmod{2^{t-\vartheta_r}},$$

und wenn $\tau_k = 2$, also jedenfalls $s - r \equiv 0 \pmod{2}$:

$$\frac{2^{v_r}}{o_1 o_2 \dots o_r} \cdot \Phi_k \equiv \psi_1 + o_{r+1} o_{r+2} \psi_2 + \dots + o_{r+1} o_{r+2} \dots o_{s-2} \psi_{\frac{s-r}{2}} \left(\text{mod } 2^{t-v_r} \right),$$

$$\psi_i \equiv \frac{2^{v_{r+2i-1}} \xi_i^2}{\varphi_{r+2i-2}} + 2 A_i \xi_i \eta_i + \frac{o_{r+2i-1} \varphi_{r+2i} - A_i^2 \varphi_{r+2i-2}}{2 \varphi_{r+2i-1}} \eta_i^2$$

wo die A_i irgendwelche ungerade Zahlen bedeuten sollen.

So oft nun $\tau_k = 1$ ist und die Zahlen $\sigma_{r-1} 2^{v_r}$ und $2^{v_s} \sigma_{s+1}$ beide durch 4 teilbar sind, kommen für den Ausdruck $f\{2\}$ zwei aus den Koeffizienten von Φ_k gebildete Einheiten ε und δ in Betracht. Indem man wiederholt die für ungerade a und α geltende Kongruenz

$$\frac{a\alpha-1}{2} \equiv \frac{a-1}{2} + \frac{\alpha-1}{2} \pmod{2}$$

anwendet, ergibt sich ε als eine Potenz von -1 mit dem Exponenten:

$$\frac{\varphi_r-1}{2} + \frac{\varphi_s-1}{2} + \sum \frac{o_{s-2u+1}+1}{2}, \quad (u=1, 2, \dots, \left[\frac{s-r}{2} \right])$$

während δ aus zwei Potenzen von -1 zusammengesetzt erscheint, von denen die eine den Exponenten:

$$\frac{\varphi_r-1}{2} \cdot \frac{\varphi_{r+1}-1}{2} + \frac{\varphi_{r+1}-1}{2} \cdot \frac{\varphi_{r+2}-1}{2} + \dots + \frac{\varphi_{s-1}-1}{2} \cdot \frac{\varphi_s-1}{2}$$

$$+ \sum_{h=r+1}^{s-1} \frac{\varphi_h-1}{2} \cdot \frac{o_h+1}{2}$$

und die andere den Exponenten:

$$\frac{\varphi_r + (-1)^{s-r}}{2} \cdot \left(\frac{\varphi_s-1}{2} + \sum \frac{o_{s-2u+1}+1}{2} \right) + \frac{\varphi_s-1}{2} \cdot \sum \frac{o_{s-2v}+1}{2}$$

$$+ \sum \frac{o_{s-2u+1}+1}{2} \cdot \frac{o_{s-2v}+1}{2} \quad \left(\begin{array}{l} u, v, u, v = 1, 2, \dots, \left[\frac{s-r}{2} \right] \\ u \leq v \end{array} \right)$$

erhält.

Die erste Potenz aus δ findet man mit Hilfe der Formeln aus *F. Q.*, p. 85 [[S. 73]] gleich

$$\left\{ \left(\frac{\varphi_{r+1}}{e_r \varphi_r} \right) (-1)^{\frac{I_r(I_r+1)}{2}} \right\} \left\{ \left(\frac{\varphi_{s-1}}{e_s \varphi_s} \right) (-1)^{\frac{I_s(I_s-1)}{2}} \right\} \cdot \prod_{h=r+1}^{s-1} \left(\frac{\varphi_h}{o_h} \right),$$

während die zweite, ebenso wie die Einheit ε sich unmittelbar durch die

Charaktere $(-1)^{\frac{\varphi_r-1}{2}}$ und $(-1)^{\frac{\varphi_s-1}{2}}$ ausdrückt.

Es verdient beachtet zu werden, daß in $f\{2\}$ die Einheiten ε und δ nur in der Verbindung $\frac{\delta + \delta^-}{2}$ auftreten, wo $\delta^- = \delta \cdot \varepsilon^{k-1}$ die zu $-\Phi_k$ gehörige Einheit δ bedeutet.

So oft ferner $\tau_k = 2$ ist, und die Zahlen $\sigma_{r-1} 2^{v_r}$ und $2^{v_s} \sigma_{s+1}$ beide

durch 4 teilbar sind, begegnet man in $f\{2\}$ einer aus den Koeffizienten von Φ_k gebildeten Einheit

$$\theta = \left(\frac{2}{o_{r+1} o_{r+3} \cdots o_{s-3} o_{s-1} \varphi_r \varphi_s} \right).$$

Wir bemerken schließlich, daß die mit $f' = \sum a_{ik} x_i x_k$ adjungierte Form $f' = \frac{(-1)^I}{d_{n-2}} \cdot \sum \frac{\partial \Delta}{\partial a_{n-i+1, n-k+1}} x'_i x'_k$ lauter Größen $f'\{q\}$ liefert, welche mit den Größen $f\{q\}$ identisch sind. Es erhellt dieses leicht aus dem Umstande, daß die Invarianten o_h' , σ_h' und die Zahlen φ_h' der Form f' die Gleichungen erfüllen:

$$o_h' = o_{n-h}, \quad \sigma_h' = \sigma_{n-h}, \quad \varphi_h' = (-1)^I \cdot \varphi_{n-h}.$$

8. Ausdruck für die Formenanzahl eines Genus.

Irgendein primitives Genus f sei definiert durch seine Invarianten I , o_h , σ_h , C , welche allen für die Existenz des Genus notwendigen Bedingungen genügen mögen. Wir bilden nach den in 5., 6. und 7. angegebenen Regeln die verschiedenen Größen $f\{q\}$, welche zu unserem Genus gehören. Wir setzen ferner

$$\mathfrak{D} = \prod_{h=1}^{n-1} o_h^{h(n-h)}$$

und

$$c_n = 2 \frac{\Gamma\left(\frac{1}{2}\right) \cdot \Gamma\left(\frac{2}{2}\right) \cdots \Gamma\left(\frac{n}{2}\right)}{\left\{ \Gamma\left(\frac{1}{2}\right) \right\}^{\frac{n(n+1)}{2}}},$$

wo $\Gamma\left(\frac{1}{2}\right) = \sqrt{\pi}$, $\Gamma(1) = 1$, $\Gamma(u+1) = u\Gamma(u)$.

Wir wollen von dem Falle absehen, wo $n=2$ und $(-1)^I o_1 = \Delta$ eine negative Quadratzahl ist, das Genus also lauter zerlegbare Formen enthält.

Schließt man diesen Fall aus, so besitzt das über alle möglichen Primzahlen $q=2, 3, 5, \dots$ erstreckte Produkt

$$(16) \quad M = c_n \cdot \frac{\sqrt{\mathfrak{D}}}{\sigma_1 \sigma_2 \cdots \sigma_{n-1}} \cdot \frac{1}{f\{2\}} \cdot \frac{1}{f\{3\}} \cdot \frac{1}{f\{5\}} \cdots \frac{1}{f\{q\}} \cdots$$

stets einen endlichen Wert.

Um dieses nachzuweisen, wollen wir in M die Größe

$$\frac{\left(1 - \frac{1}{q^2}\right) \left(1 - \frac{1}{q^4}\right) \cdots \left(1 - \frac{1}{q^{\left[\frac{n-1}{2} \right]}}\right)}{f\{q\}} = E_q$$

als allgemeines Glied einführen. Solches kann leicht geschehen, indem man mit der Identität

$$1 = S_2 S_4 \dots S_{2^{\left[\frac{n-1}{2}\right]}} \cdot \prod_q \left(1 - \frac{1}{q^2}\right) \left(1 - \frac{1}{q^4}\right) \dots \left(1 - \frac{1}{q^{2^{\left[\frac{n-1}{2}\right]}}}\right)$$

multipliziert, wo S_{2^k} die Summe $\sum_{z=1}^{\infty} \frac{1}{z^{2^k}}$ bedeutet. Man weiß, daß diese Summe den Wert $\frac{1}{2} B_k \cdot \frac{(2\pi)^{2k}}{(2k)!}$ hat, falls unter B_k die k^{te} Bernoullische Zahl verstanden wird.

Für jede nicht in 2Δ enthaltene Primzahl p kommt, wenn $n \equiv 1 \pmod{2}$:

$$E_p = 1,$$

und wenn $n \equiv 0 \pmod{2}$:

$$E_p = \left\{ 1 - \left(\frac{(-1)^{\frac{n}{2}} \Delta}{p} \right) p^{-\frac{n}{2}} \right\}^{-1}, \quad \left(\frac{\Delta}{p} \right) = \left(\frac{(-1)^{I \cdot \mathfrak{D}}}{p} \right)$$

Sind also die nicht in Δ aufgehenden, ungeraden Primzahlen, ihrer Größe nach geordnet: p, p', p'' usw., so wird das unendliche Produkt:

$$E_p E_{p'} E_{p''} \dots,$$

je nachdem $n \equiv 1 \pmod{2}$ oder $n \equiv 0 \pmod{2}$, gleich 1 oder gleich der Summe:

$$\sum \left(\frac{(-1)^{\frac{n}{2}} \Delta}{m} \right) \frac{1}{m^{\frac{n}{2}}},$$

wo m alle positiven und zu 2Δ relativ primen ganzen Zahlen durchläuft. Zur Bezeichnung dieser Dirichletschen Summe mag das Symbol

$$D_{\frac{n}{2}} \left[(-1)^{\frac{n}{2}} \Delta \right]$$

dienen.

Man setze nun:

$$c_n \cdot S_2 S_4 \dots S_{2^{\left[\frac{n-1}{2}\right]}} = c_n,$$

d. i., wenn $n \equiv 1 \pmod{2}$:

$$c_n = \left(\frac{1}{2} \right)^{\frac{n-3}{2}} \cdot B_1 B_2 \dots B_{\frac{n-1}{2}} \cdot \frac{1}{1 \cdot 2 \dots \frac{n-1}{2}},$$

und wenn $n \equiv 0 \pmod{2}$:

$$c_n = \left(\frac{1}{2} \right)^{\frac{n-4}{2}} \cdot B_1 B_2 \dots B_{\frac{n-2}{2}} \cdot \frac{1}{\pi^{\frac{n}{2}}},$$

und lasse ferner $\mathfrak{d}$ alle Primzahlen aus $2\mathfrak{D}$ durchlaufen. Dann können wir endlich schreiben, wenn $n \equiv 1 \pmod{2}$:

$$(17) \quad M = c_n \cdot \frac{\sqrt{\mathfrak{D}}}{\sigma_1 \sigma_2 \dots \sigma_{n-1}} \cdot \prod_{\partial} E_{\partial},$$

und wenn $n \equiv 0 \pmod{2}$:

$$(17) \quad M = c_n \cdot \frac{\sqrt{\mathfrak{D}}}{\sigma_1 \sigma_2 \dots \sigma_{n-1}} \cdot \prod_{\partial} E_{\partial} \cdot D_{\frac{n}{2}} \left[(-1)^{\frac{n}{2}-I} \mathfrak{D} \right].$$

Diese Ausdrücke zeigen in der Tat, daß M einen endlichen und positiven Wert annimmt.

Für ein Genus von binären zerlegbaren Formen gewinnt man ein ähnliches konvergentes Produkt M , indem man $\frac{1 - \frac{1}{q}}{f\{q\}}$ an die Stelle von $\frac{1}{f\{q\}}$ treten läßt.

Wir behaupten nun:

Die Formenanzahl unseres Genus besitzt den Ausdruck:

$$M \cdot \Omega,$$

wo M das angegebene Produkt und Ω eine positive unendliche GröÙe bedeutet, die nur von n und I und von der Anzahl der Darstellungen der Zahl $()$ durch die Formen des Genus abhängt.

In den Fällen eines definiten Genus ($I=0$ oder $I=n$) ist insbesondere dieses Ω gleich der Anzahl aller ganzzahligen n -reihigen Substitutionen von der Determinante 1, und mithin M gleich der über alle Klassen Cl des Genus erstreckten Summe $\sum \frac{1}{t(Cl)}$.

Wir werden uns begnügen, das vorstehende Resultat für die Fälle definiten (und zwar positiver) Genera zu erweisen. Dabei werden wir von den Dirichletschen Methoden*) Gebrauch machen, und in den Fällen $n > 2$ einen Schluß von $n-1$ auf n zu Hilfe nehmen.

Zweiter Teil.

9. Das Maß eines positiven Genus dargestellt durch einen gewissen Grenzwert.

Ein positives Genus G von $n (\geq 2)$ Variablen sei definiert durch seine Ordnung O :

$$d_0, \quad \left(\begin{matrix} \sigma_1, \sigma_2, \dots, \sigma_{n-2}, \sigma_{n-1} \\ \sigma_1, \sigma_2, \dots, \sigma_{n-2}, \sigma_{n-1} \end{matrix} \right), \quad I=0$$

und seine Charaktere C . Wir setzen voraus, daß dieselben den für die Existenz des Genus notwendigen Bedingungen genügen. d_0 sei gleich 1, das Genus also primitiv.

*) Dirichlet, Recherches sur diverses applications de l'analyse infinitésimale à la théorie des nombres. (Crelles Journal, Bd. 19, und Werke, Bd. I.)

Es sei Δ die Determinante des Genus, und R eine beliebige zu 2Δ relativ prime ganze Zahl. Durch die Kenntnis der Invarianten O und C sind wir befähigt, die Reste unseres Genus für einen jeden beliebigen Modul hinzuschreiben. Insbesondere können wir also irgendeinen Hauptrest φ in bezug auf den Modul $\sigma_1 \cdot 8\Delta R$ angeben. Der erste Koeffizient von φ heiße $\sigma_1 \alpha$. Die Zahl α ist dann sicher relativ prim zu $8\Delta R$.

Wir richten unser Augenmerk auf den quadratischen Charakter von α in bezug auf den Modul $8\Delta R$. Enthält Δ im ganzen δ ungerade Primzahlen ∂ , und R im ganzen r ungerade Primzahlen r , so wird dieser Charakter durch die Gesamtheit der folgenden $2 + \delta + r$ Symbole definiert:

$$C(\alpha) \quad (-1)^{\frac{\alpha-1}{2}}, \quad \left(\frac{2}{\alpha}\right), \quad \left(\frac{\alpha}{\partial}\right), \quad \left(\frac{\alpha}{r}\right).$$

Diese Symbole können zum Teil Charaktere des Genus vorstellen, zum Teil können sie bei anderer Wahl des Restes φ andere Werte erlangen. *)

*) Man überzeugt sich leicht, daß in dieser Beziehung die nachstehenden Sätze gelten:

Eine jede Einheit $\left(\frac{\alpha}{r}\right)$ kann sowohl gleich $+1$ wie gleich -1 ausfallen.

Eine Einheit $\left(\frac{\alpha}{\partial}\right)$ hat einen festen Wert, wenn $\sigma_1 \equiv 0 \pmod{\partial}$, also φ von dem Typus $\sigma_1 \alpha \xi^2 \pmod{\partial}$ ist. Sonst kann dieselbe beide Werte ± 1 annehmen.

Was die Einheiten $(-1)^{\frac{\alpha-1}{2}}$ und $\left(\frac{2}{\alpha}\right)$ anlangt, so unterliegen dieselben keiner Beschränkung, wenn $\sigma_1 = 2$ ist. Ebenso im allgemeinen, wenn $\sigma_1 = 1$ ist; nur bestehen hier die folgenden Ausnahmefälle:

1. Ist $\varphi \equiv \alpha \xi^2 \pmod{8}$, so sind beide Einheiten $(-1)^{\frac{\alpha-1}{2}}$ und $\left(\frac{2}{\alpha}\right)$ Charaktere.
 2. Ist $\varphi \equiv \alpha \xi^2 \pmod{4}$, oder $\varphi \equiv \alpha \xi^2 + \beta \eta^2 \pmod{4}$ und $\alpha \equiv \beta \pmod{4}$, oder $\varphi \equiv \alpha \xi^2 + \beta \eta^2 + \gamma \zeta^2 \pmod{4}$ und $\alpha \equiv \beta \equiv \gamma \pmod{4}$, so hat die Einheit $(-1)^{\frac{\alpha-1}{2}}$ einen festen Wert.

3. Ist $\varphi \equiv \alpha \xi^2 + 2\beta \eta^2 \pmod{8}$ und setzt man $(-1)^{\frac{\alpha\beta+1}{2}} = \varepsilon$, so können $(-1)^{\frac{\alpha-1}{2}}$ und $\left(\frac{2}{\alpha}\right)$ sich nur mit φ so ändern, daß $\varepsilon^{\frac{\alpha-1}{2}} \cdot \left(\frac{2}{\alpha}\right)$ fest bleibt.

4. Wenn $\varphi \equiv \alpha \xi^2 + 2(\beta \eta^2 + \gamma \zeta^2) \pmod{8}$, so sind für die Einheiten $(-1)^{\frac{\alpha-1}{2}}$ und $\left(\frac{2}{\alpha}\right)$ nur drei von den vier Systemen $\pm 1, \pm 1$ zulässig. Ist nämlich $\beta \equiv -\gamma \pmod{4}$, so kann der Fall nicht eintreten, daß $(-1)^{\frac{\alpha-1}{2}}$ ungeändert bleibt, während $\left(\frac{2}{\alpha}\right)$ in $-\left(\frac{2}{\alpha}\right)$ übergeht, und hat man $\beta \equiv \gamma \pmod{4}$, $\theta \equiv \pm 1$, so ist der Fall ausgeschlossen, daß $(-1)^{\frac{\alpha-1}{2}}$ ins Gegenteil umschlägt, während $\left(\frac{2}{\alpha}\right)$ sich in $\theta \cdot \left(\frac{2}{\alpha}\right)$ verwandelt.

Wie in 6. bezeichnen wir mit κ den Index der ersten von den n Zahlen $o_1, o_2, \dots, o_{n-1}, o_n (= 0)$, welche gerade ausfällt.

Wir denken uns in jeder überhaupt existierenden Klasse des Genus je eine Form ausgesucht, welche nach dem Modul $\sigma_1 \cdot 8\Delta R$ den Rest φ läßt. Eine beliebige der so gewonnenen Formen sei f .

Wir bestimmen für die Variablen von f alle Systeme $\xi_1, \xi_2, \dots, \xi_n$, welche dem Ausdrucke $f(\xi_i)$ einen Wert $\sigma_1 m$ erteilen, wo m prim zu $8\Delta R$ ist und den Gleichungen

$$C(m) = C(\alpha)$$

genügt, welche aber dabei, falls $\kappa > 1$ ist, nicht die Bedingungen

$$(7) \quad \xi_1 \equiv \xi_2 \equiv \dots \equiv \xi_\kappa \equiv \sigma_1 \pmod{2}$$

erfüllen. Über alle diese Wertsysteme $\xi_i (\neq 0, 0, \dots, 0)$ erstrecken wir alsdann die Summe

$$\equiv \varrho \sum \frac{1}{\left\{ \frac{1}{\sigma_1} f(\xi_i) \right\}^{\frac{n}{2}(1+\varphi)}},$$

und wir wollen den Grenzwert ermitteln, welchen diese Summe für ein positives, unendlich abnehmendes ϱ erreicht.

Die definierten Wertsysteme ξ_i werden in einer gewissen Anzahl A von arithmetischen Progressionen

$$\xi_1 = 8\Delta R \cdot X_1 + v_1, \quad \xi_2 = 8\Delta R \cdot X_2 + v_2, \quad \dots, \quad \xi_n = 8\Delta R \cdot X_n + v_n$$

enthalten sein. Man bezeichne mit $2^{8(n-1)} \cdot A_2$, wenn $\kappa > 1$, die Anzahl aller derjenigen Lösungen $\xi_i \pmod{8}$ von

$$\frac{1}{\sigma_1} \varphi(\xi_i) \equiv \alpha \pmod{8},$$

welche nicht zugleich den Bedingungen (7) genügen, und wenn $\kappa = 1$, die Anzahl aller möglichen Lösungen dieser Kongruenz; ferner mit $2^{n-1} \cdot A_p$ die Anzahl aller Lösungen von

$$\varphi(\xi_i) \equiv \sigma_1 \alpha \pmod{p},$$

wenn p eine der ungeraden Primzahlen aus ΔR bedeutet. In den Ausdrücken von A_2 und A_p erscheint α , wie wir wissen, nur in den Einheiten $C(\alpha)$. Dieser Umstand läßt erkennen, daß die Anzahl A den Wert hat:

$$A = (8\Delta R)^n \cdot \frac{1}{2} \prod_q \left\{ \frac{1}{2} \left(1 - \frac{1}{q} \right) A_q \right\}, \quad (q = 2, \varrho, r)$$

wo q die $1 + \mathfrak{b} + \mathfrak{r}$ Primzahlen von $8\Delta R$ durchläuft.

Setzt man für die ξ_i zunächst nur alle diejenigen Systeme, welche in einer der A Progressionen vorkommen, so ergibt sich der Grenzwert der entstehenden Summe für ein $\varrho = +0$, nach *F. Q.*, p. 148 [[S. 121]], gleich

$$e_n \cdot \frac{(8\Delta R)^{-n}}{\sqrt{\frac{\Delta}{\sigma_1^n}}},$$

wo

$$e_n = \frac{\left\{ \Gamma\left(\frac{1}{2}\right) \right\}^n}{\Gamma\left(1 + \frac{n}{2}\right)} = \pi^{\left[\frac{n}{2}\right]} \cdot \frac{2^{\left[\frac{n+1}{2}\right]}}{n(n-2) \cdots \left(n-2\left[\frac{n-1}{2}\right]\right)}.$$

Für die ganze Summe Ξ wird daher

$$\lim_{\varrho=0} (\Xi) = e_n \cdot \frac{(8\Delta R)_1}{2^{2+\frac{1}{2}+\tau}} \cdot \frac{\sigma_1^{\frac{n}{2}}}{\sqrt{\Delta}} \cdot \prod_q A_q, \quad (q=2, 3, \tau)$$

wo $(8\Delta R)_1$ das über alle Primzahlen q von $8\Delta R$ erstreckte Produkt $\prod \left(1 - \frac{1}{q}\right)$ bedeutet.

Eine Summe X , von ähnlicher Beschaffenheit wie Ξ , mag jetzt nur alle solchen Systeme $\xi_i = x_i$ umfassen, in welchen die n Zahlen $\xi_1, \xi_2, \dots, \xi_n$ keinen gemeinschaftlichen Teiler haben. Offenbar entstehen alle möglichen Systeme ξ_i , wenn wir diese besonderen Systeme x_i mit allen positiven und zu $8\Delta R$ relativ primen Zahlen z multiplizieren. Man findet daher

$$\Xi = X \cdot \sum \frac{1}{z^{n(1+\varrho)}},$$

und in der Grenze für $\varrho = 0$:

$$\lim \left(\frac{\Xi}{X} \right) = \sum \frac{1}{z^n}.$$

Die hier auftretende Summe hat bekanntlich den Wert:

$$(8\Delta R)_n \cdot S_n,$$

wenn S_n die Summe $1 + \frac{1}{2^n} + \frac{1}{3^n} + \dots$ und $(8\Delta R)_n$ das über alle Primzahlen q von $8\Delta R$ erstreckte Produkt $\prod \left(1 - \frac{1}{q^n}\right)$ ausdrückt.

Wir bemerken noch, daß die Summe X sich in $t(f)$ untereinander identische Summen X_0 zerlegen läßt. Dabei ist unter $t(f)$, wie in 1., die Anzahl aller Substitutionen von der Determinante 1 verstanden, welche die Form f in sich selbst transformieren. (Solche Summen X_0 haben dann auch in Fällen indefiniter Formen einen Sinn.)

Wir bilden endlich die Doppelsumme:

$$S = \varrho \cdot \sum t(f) \cdot \sum \frac{1}{\left\{ \frac{f(x_i)}{\sigma_1} \right\}^{\frac{n}{2}(1+\varrho)}},$$

erstreckt, einmal: über alle die inäquivalenten Formen $f(\equiv \varphi, \text{ mod } \sigma_1 \cdot 8\Delta R)$, die wir oben als Repräsentanten der einzelnen Klassen des Genus auf-

gestellt hatten; und dann, für jede dieser Formen f : über alle solchen Systeme $x_1, x_2, \dots, x_n$ ohne gemeinsamen Teiler, welche einer Kongruenz

$$\frac{f(x_i)}{\sigma_1} \equiv \alpha z^2 \pmod{8\Delta R}$$

genügen, wo z zu $8\Delta R$ relativ prim ist, und welche dabei, falls $\kappa > 1$, nicht alle Bedingungen

$$(7) \quad x_1 \equiv x_2 \equiv \dots \equiv x_\kappa \equiv \sigma_1 \pmod{2}$$

erfüllen.

Der Grenzwert l der früheren Summe X hatte sich als unabhängig von der speziellen Form f erwiesen. Infolgedessen muß der Grenzwert L dieser Doppelsumme gleich $l \propto \sum \frac{1}{t(f)}$ sein. Durch die hier erscheinende einfache Summe ist aber das Maß M unseres Genus dargestellt; man erhält demnach:

$$L = e_n \cdot \frac{(8\Delta R)_1}{2^{2+b+\tau}} \cdot \frac{1}{(8\Delta R)_n \cdot S_n} \cdot \frac{\sigma_1^{\frac{n}{2}}}{\sqrt{\Delta}} \cdot \prod_q A_q \cdot M. \quad (q=2, \partial, r)$$

Wir werden jetzt für L einen zweiten Ausdruck ableiten, und durch Vergleichung der beiden Ausdrücke werden wir dann die in 8. aufgestellten Formeln als richtig erkennen.

10. Bestimmung desselben Grenzwertes auf anderem Wege.

Wir haben soeben den Grenzwert L gefunden, indem wir uns die Summe S erst nach den einzelnen Formen f , und dann nach der numerischen Größe der Zahlen $\frac{f(x_i)}{\sigma_1}$ geordnet dachten. Nun handelt es sich in S um lauter positive Glieder; wir müssen daher zu demselben Grenzwert L gelangen, wenn wir die Summe S direkt nach der Größe der Zahlen $\frac{1}{\sigma_1} f(x_i) = m$ anordnen. Durch ein solches Arrangement entsteht für S zunächst ein Ausdruck von der Gestalt:

$$\varrho \sum_m \frac{M(m)}{\frac{n}{2}(1+\varrho)},$$

wo die Summation alle positiven Zahlen m betrifft, die zu $8\Delta R$ relativ prim sind und den Gleichungen $C(m) = C(\alpha)$ genügen.

Für jede dieser Zahlen m hat die Größe $M(m)$ folgende Bedeutung. Es bezeichne $m(f)$, wie oft eine bestimmte Form f die Zahl $\sigma_1 m$ mit Hilfe von Systemen x_i darzustellen vermag, welche keinen gemeinsamen Teiler > 1 haben, und außerdem, falls $\kappa > 1$, nicht alle Bedingungen (7) erfüllen. Dann ist:

$$M(m) = \sum \frac{m(f)}{t(f)},$$

wo die Summe sich über alle die Formen f erstreckt. Die Größe $M(m)$ bildet also, wie wir uns nach *F. Q.*, p. 142 [[S. 116]] ausdrücken, das Maß für alle die definierten Darstellungen x_i der Zahl $\sigma_1 m$ durch die verschiedenen Formen f des Genus G .

Treten in der Zahl m im ganzen μ ungerade Primzahlen $p_1, p_2, \dots, p_\mu$ auf, so erscheint, nach *F. Q.*, p. 143 [[S. 117]], dieses Maß $M(m)$ als das 2^μ -fache von dem Maße eines bestimmten positiven Genus $G(m)$ von Formen mit $n - 1$ Variablen, welches enge mit dem Genus G zusammenhängt. Von den Sätzen, welche diesen Zusammenhang feststellen, wollen wir hier soviel anführen, als für die Bestimmung der Größe $M(m)$ von Wichtigkeit ist (vgl. *F. Q.*, art. XVIII [[S. 102—113]]).

1. Es sei zunächst $n = 2$, in welchem Falle Δ und σ_1 identisch sind. Je nach der Beschaffenheit der Zahl m bieten sich zwei Möglichkeiten dar.

Entweder ist für die Zahl m die quadratische Kongruenz

$$-\Delta \equiv z^2 \pmod{\sigma_1 m}$$

nicht lösbar. In diesem Falle existiert auch das Genus $G(m)$ nicht, und man hat sein Maß gleich 0 zu setzen.

Oder diese Kongruenz ist lösbar und besitzt 2^μ Lösungen $z \pmod{\sigma_1 m}$. Alsdann wird das Genus $G(m)$ von der einen Form $g = \xi^2$ gebildet und liefert das Maß 1.

In beiden Fällen kann man schreiben:

$$M(m) = \left\{ 1 + \left(\frac{-\Delta R^2}{p_1} \right) \right\} \left\{ 1 + \left(\frac{-\Delta R^2}{p_2} \right) \right\} \cdots \left\{ 1 + \left(\frac{-\Delta R^2}{p_\mu} \right) \right\}.$$

Der zweite Fall ereignet sich beispielsweise, wenn man für $\sigma_1 m$ den ersten Koeffizienten einer der Formen f wählt, was man tun darf. Denn nach Voraussetzung ist ein solcher Koeffizient $\equiv \sigma_1 \alpha \pmod{\sigma_1 \cdot 8\Delta R}$. Man hat dann jedenfalls $\left(\frac{-\Delta}{m} \right) = 1$, worin eine Bedingung für die Einheiten $C(m) = C(\alpha)$ liegt.

2. Ist $n > 2$, so erkennt man zunächst, daß die Invarianten von $G(m)$ durch die Zahl m [$C(m) = C(\alpha)$] und die Invarianten von G stets in solcher Weise ausgedrückt sind, daß das Genus $G(m)$ wirklich existiert. Wir haben bereits bemerkt, daß dieses Genus sich als positiv erweist. Man findet es auch primitiv. Seine Invarianten o und σ erlangen die Werte:

$$\left(\begin{array}{c} \sigma_2, \sigma_3, \dots, \sigma_{n-1} \\ \sigma_1 m \cdot o_2, o_3, \dots, o_{n-1} \end{array} \right).$$

Es sei $\Delta^1 = \prod_{h=2}^{n-1} o_h^{n-h}$ und $\mathfrak{D}^1 = \prod_{h=2}^{n-1} o_h^{(h-1)(n-h)}$. Ferner fallen seine Cha-

raktere derart aus, daß für jede seiner Formen $g = \sum_1^{n-1} c_{ik} \xi_i \xi_k$ die $\frac{n(n-1)}{2}$

Kongruenzen:

$$(18) \quad -o_1 c_{ik} \equiv z_i z_k \pmod{\sigma_1 m}$$

je 2^μ Lösungen zulassen.

Wir nehmen nun an, die Formeln aus 8. seien bereits für den Fall $n-1$ erwiesen, und wir wollen dieselben benutzen, um das Maß von $G(m)$ aufzustellen. Wir bilden zu dem Behufe für irgendeine Form g dieses Genus alle Größen $g\{q\}$, und wir schreiben:

$$\frac{\left(1 - \frac{1}{q^2}\right) \left(1 - \frac{1}{q^4}\right) \cdots \left(1 - \frac{1}{q^{2\left[\frac{n}{2}\right] - 2}}\right)}{g\{q\}} = E_q^1.$$

Für das Maß von $G(m)$ erhalten wir dann einen Ausdruck:

$$c_{n-1} \cdot \frac{\sqrt{\mathfrak{D}^1} \cdot \sigma_1^{\frac{n-2}{2}} m^{\frac{n-2}{2}}}{\sigma_2 \sigma_3 \cdots \sigma_{n-1}} \cdot E_2^1 E_3^1 E_5^1 \cdots E_q^1 \cdots,$$

wo c_{n-1} eine Konstante bedeutet; und die Größe $M(m)$ ergibt sich gleich diesem Ausdrucke, multipliziert mit 2^μ . Es sind jetzt die Größen E_q^1 zu ermitteln.

1) Ist zunächst q irgendeine ungerade Primzahl, die weder in ΔR , noch in m aufgeht, so kommt nach 8., wenn $n \equiv 0 \pmod{2}$:

$$E_q^1 = 1,$$

und wenn $n \equiv 1 \pmod{2}$:

$$E_q^1 = \left[1 - \left(\frac{(-1)^{\frac{n-1}{2}} \sigma_1 m \Delta^1}{q} \right)^{\frac{1}{\frac{n-1}{2}}} \right]^{-1}.$$

In dem letzteren Falle hat man zugleich: $\left(\frac{\Delta^1}{q}\right) = \left(\frac{\Delta R^1}{q}\right)$. Bildet man das Produkt $\prod E_q^1$ über alle nicht in $2\Delta Rm$ aufgehenden Primzahlen q in ihrer natürlichen Reihenfolge, so kann man für dasselbe, je nachdem $n \equiv 0 \pmod{2}$ oder $n \equiv 1 \pmod{2}$ ist, entweder 1 oder die Summe

$$D_{\frac{n-1}{2}} \left[(-1)^{\frac{n-1}{2}} \sigma_1 m \Delta R^2 \right] \text{ setzen.}$$

2) Ist weiter $q = p$ eine der μ ungeraden Primzahlen aus m , und p^d die höchste Potenz dieser Primzahl, welche m teilt, so besitzt das Genus $G(m)$ Reste vom Typus:

$$g \equiv c \xi^2 + p^d (c_1 \xi_1^2 + \cdots + c_{n-2} \xi_{n-2}^2) \pmod{p^d}, \quad (t > d)$$

wo die Größen $c, c_1, \dots, c_{n-2}$ sämtlich zu p prim sind. Aus den Kongruenzen (18) erschließt man das Bestehen einer Kongruenz:

$$-o_1 c \equiv z^2 \pmod{p^d};$$

man findet ferner:

$$cc_1 \dots c_{n-2} \equiv \left(\frac{\sigma_1 m}{p^d}\right)^{n-2} \cdot \Delta^1 \pmod{p^{t-d}}.$$

Im Falle eines $n \equiv 0 \pmod{2}$, wo zugleich $\left(\frac{o_1 \Delta^1}{p}\right) = \left(\frac{\Delta}{p}\right)$, kommt also

$$\left(\frac{c_1 \dots c_{n-2}}{p}\right) = \left(\frac{-\Delta}{p}\right).$$

Die Formeln aus 5. und 8. geben in diesem Falle:

$$E_p^1 = \frac{1}{2} \left[1 + \left(\frac{(-1)^{\frac{n-2}{2}} c_1 \dots c_{n-2}}{p} \right) \frac{1}{p^{\frac{n-2}{2}}} \right] = \frac{1}{2} \left\{ 1 + \left(\frac{(-1)^{\frac{n}{2}} \Delta R^2}{p} \right) \frac{1}{p^{\frac{n-2}{2}}} \right\},$$

dagegen, wenn $n \equiv 1 \pmod{2}$:

$$E_p^1 = \frac{1}{2}.$$

3) Ist endlich q eine der Primzahlen aus $8\Delta R$, so besitzt das gegebene Genus G für einen jeden Modul q^t Hauptreste f_1 mit einem ersten Koeffizienten $\sigma_1 m$. Aus diesen Hauptresten entspringen, nach den Sätzen aus *F. Q.*, art. XVIII [[S. 102—113]], in einfacher Weise Hauptreste des Genus $G(m)$.

Bedeutet q zunächst eine der ungeraden Primzahlen ϱ, r , so hat ein f_1 den Typus:

$$f_1 \equiv \sigma_1 m \xi^2 + \frac{o_1 f^{(1)}}{\sigma_1 m} \pmod{q^t}.$$

Der hier auftretende Rest $f^{(1)} \pmod{q^{t-\omega_1}}$ bildet dann einen Rest des Genus $G(m)$.

Ist $q = 2$, so müssen die Fälle $\sigma_1 = 1$ und $\sigma_1 = 2$ unterschieden werden.

Im ersteren Falle hat ein f_1 den Typus:

$$f_1 \equiv m \xi^2 + \frac{o_1 f^{(1)}}{m} \pmod{2^t},$$

wo $f^{(1)}$ einen in bezug auf 2 primitiven Rest vorstellt, welcher im Falle $o_1 \equiv 1 \pmod{2}$ eine erste Invariante σ gleich 1 liefert. In diesem $f^{(1)} \pmod{2^{t-\omega_1}}$ finden wir einen Rest von $G(m)$.

Im zweiten Falle ($\sigma_1 = 2$) hat f_1 den Typus:

$$f_1 \equiv 2(m \xi^2 + A \xi \tilde{\xi} + \tilde{m} \tilde{\xi}^2) + \frac{o_1 o_2 f^{(2)}}{m} \pmod{2^t},$$

wo A ungerade und $f^{(2)}$ primitiv in bezug auf 2 ist, und man gewinnt in

$$f^{(1)} \equiv \frac{(4m\tilde{m} - A^2)}{o_1} \eta^2 + 2o_2 f^{(2)} \pmod{2^{t+1}}$$

einen Rest des Genus $G(m)$.

In allen Fällen ergibt sich nun, nach 5. und 6., wenn t groß genug gewählt ist, die Beziehung:

$$f\{q\} = A_q \cdot f^{(1)}\{q\},$$

wobei A_q mit der in 9. auf diese Weise bezeichneten Größe identisch ist. Für die Ausdrücke E_q und E_q^1 folgt hieraus, wenn $n \equiv 0 \pmod{2}$ und > 2 , die Gleichung:

$$E_q^1 = A_q E_q,$$

und wenn $n \equiv 1 \pmod{2}$:

$$E_q^1 = \frac{A_q E_q}{1 - \frac{1}{q^{n-1}}}.$$

Im Falle $n = 2$ findet man in ähnlicher Weise: $f\{q\} = A_q$, also, da hier $E_q = \frac{1}{f\{q\}}$ ist: $A_q E_q = 1$.

Fassen wir alles Vorhergehende zusammen, so können wir die Größe $M(m)$, wenn $n > 2$ ist, in folgender Form schreiben:

$$c_{n-1} \cdot \frac{\sqrt{\mathfrak{D}}^{\frac{n-2}{2}} \sigma_1^{\frac{n-2}{2}} m^{\frac{n-2}{2}}}{\sigma_2 \sigma_3 \dots \sigma_{n-1}} \cdot \prod_q A_q E_q \cdot (m), \quad (q = 2, \vartheta, r)$$

wobei im Falle eines geraden n :

$$(m) = \prod_p \left\{ 1 + \left(\frac{(-1)^{\frac{n}{2}} \Delta R^2}{p} \right)^{\frac{1}{\frac{n-2}{2}}} \right\}, \quad (p = p_1, p_2, \dots, p_\mu)$$

und im Falle eines ungeraden n :

$$(m) = \frac{1}{(8 \Delta R)_{n-1}} \cdot D_{\frac{n-1}{2}} [(-1)^{\frac{n-1}{2}} \sigma_1 m \Delta R^2].$$

Dieser Ausdruck bleibt nun auch für $n = 2$ gültig, wenn $c_1 = 1$ genommen wird.

Wir vergleichen jetzt den früher gefundenen Ausdruck von L mit dem Grenzwerte $\lim \left(\varrho \sum_{m^{\frac{n}{2}(1+\varrho)}} \frac{M(m)}{m^{\frac{n}{2}(1+\varrho)}} \right)$. Dadurch erhalten wir eine Beziehung für das Maß M des gegebenen Genus. In dieselbe setzen wir

$$M = c_n \cdot \frac{\sqrt{\mathfrak{D}}}{\sigma_1 \sigma_2 \dots \sigma_{n-1}} \cdot \prod_q E_q \cdot D_R \cdot M_0, \quad (q = 2, \vartheta, r)$$

wo $D_R = 1$ sei für ein ungerades n , und gleich $D_{\frac{n}{2}} [(-1)^{\frac{n}{2}} \Delta R^2]$ für ein gerades n , und wo $\mathfrak{D} = \Delta \cdot \mathfrak{D}^1$ und

$$\begin{aligned} c_n &= 2 \frac{c_1}{e_2} (n=2); &= \frac{c_{n-1}}{e_n} \cdot \frac{2}{n} (n \equiv 0, \pmod{2}; n > 2); \\ &= \frac{c_{n-1}}{e_n} \cdot \frac{2}{n} \cdot S_{n-1} (n \equiv 1, \pmod{2}) \end{aligned}$$

die Größen aus 8. bedeuten sollen. Die Größe M_0 erweist sich dann als unabhängig von der Zahl R , und es wird $M_0 = 1$ sein müssen, damit die in 8. für M aufgestellten Ausdrücke in Wirklichkeit gelten.

Schreibt man noch $\frac{2}{n}\varrho$ für ϱ , so lauten die Endformeln: wenn $n = 2$,

$$(19) \quad 2 \cdot \frac{(8\Delta R)_1}{2^{2+b+r}} \cdot \frac{D_1[-\Delta R^2]}{(8\Delta R)_2 \cdot S_2} \cdot M_0 = \text{Lim } \varrho \sum \frac{1}{m^{1+\varrho}} \prod_p \left[1 + \left(\frac{-\Delta R^2}{p} \right) \right];$$

wenn $n \equiv 0 \pmod{2}$ und > 2 ,

$$(20) \quad \frac{(8\Delta R)_1}{2^{2+b+r}} \cdot \frac{D_{\frac{n}{2}}[(-1)^{\frac{n}{2}} \Delta R^2]}{(8\Delta R)_n \cdot S_n} \cdot M_0 \\ = \text{Lim} \left\{ \varrho \sum \frac{1}{m^{1+\varrho}} \prod_p \left[1 + \left(\frac{(-1)^{\frac{n}{2}} \Delta R^2}{p} \right) \frac{1}{p^{\frac{n-2}{2}}} \right] \right\};$$

wenn $n \equiv 1 \pmod{2}$,

$$(21) \quad \frac{(8\Delta R)_1}{2^{2+b+r}} \cdot \frac{(8\Delta R)_{n-1} \cdot S_{n-1}}{(8\Delta R)_n \cdot S_n} \cdot M_0 \\ = \text{Lim} \left\{ \varrho \sum \frac{1}{m^{1+\varrho}} D_{\frac{n-1}{2}}[(-1)^{\frac{n-1}{2}} \sigma_1 m \Delta R^2] \right\}.$$

Dabei durchläuft m , wie erinnert werden mag, alle positiven und zu $8\Delta R$ relativ primen ganzen Zahlen, für welche die $2 + b + r$ Einheiten

$$C(m) \quad (-1)^{\frac{m-1}{2}}, \quad \left(\frac{2}{m} \right), \quad \left(\frac{m}{\partial} \right), \quad \left(\frac{m}{r} \right)$$

gewisse feste Werte annehmen, die für ein $n = 2$ jedenfalls der Bedingung $\left(\frac{-\Delta R^2}{m} \right) = 1$ genügen.

Diese Zahlen m bilden offenbar die Individuen von $(8\Delta R)^n \cdot \frac{(8\Delta R)_1}{2^{2+b+r}}$ arithmetischen Progressionen $8\Delta R \cdot U + m_0$, ($U = 0, 1, \dots, \infty$) von der Differenz $8\Delta R$. Infolgedessen muß nach Dirichlet die Gleichung bestehen:

$$(22) \quad \frac{(8\Delta R)_1}{2^{2+b+r}} = \text{Lim} \left(\varrho \sum \frac{1}{m^{1+\varrho}} \right).$$

Von derselben werden wir sofort Gebrauch machen, um in allen Fällen das Resultat $M_0 = 1$ abzuleiten.

11. Beweis, daß $M_0 = 1$ ist.

Wir schicken die folgende Betrachtung voraus:

Es sei eine ganze positive oder negative Zahl N teilbar durch alle ungeraden Primzahlen, die unter einer gewissen Grenze $G + 1$ liegen;

ferner soll p die sämtlichen Primzahlen irgendeiner zu $2N$ relativ primen Zahl m durchlaufen; endlich sei $\nu > 1$ und $\gamma = \frac{1}{(\nu-1)G^{\nu-1}}$; dann gelten die Ungleichungen:

$$1 - \gamma < D_\nu[N] < 1 + \gamma; \quad 1 < (2N)_\nu \cdot S_\nu < 1 + \gamma;$$

$$1 - \gamma < \prod \left[1 + \left(\frac{N}{p} \right) \frac{1}{p^\nu} \right] = II_\nu < 1 + \gamma.$$

In der Tat, man erhält zunächst

$$D_\nu[N] = \sum \left(\frac{N}{m} \right) \frac{1}{m^\nu},$$

wo die Summation sich auf alle zu $2N$ relativ primen und positiven Zahlen m bezieht. Die kleinste dieser Zahlen m , welche von 1 verschieden ist, besitzt mindestens den Wert $G+1$. Die vorstehende Summe liegt daher zwischen den beiden Größen:

$$1 \pm \left(\frac{1}{(G+1)^\nu} + \frac{1}{(G+2)^\nu} + \dots \right).$$

Da nun

$$\frac{1}{(G+k)^\nu} < \int_{G+k-1}^{G+k} \frac{dx}{x^\nu},$$

so folgt um so mehr:

$$1 - \int_G^\infty \frac{dx}{x^\nu} < D_\nu[N] < 1 + \int_G^\infty \frac{dx}{x^\nu}.$$

In derselben Weise ergeben sich die Ungleichungen für die Größe $(2N)_\nu \cdot S_\nu$, welche den Wert der über alle Zahlen m erstreckten Summe $\sum \frac{1}{m^\nu}$ ausdrückt.

Was endlich das Produkt II_ν anlangt, so kommt zunächst

$$II_\nu \leq \prod (1 + p^{-\nu}) < 1 + [(G+1)^{-\nu} + (G+2)^{-\nu} + \dots] < 1 + \gamma,$$

und dann

$$II_\nu \geq \prod (1 - p^{-\nu}) > \frac{1}{1 + [(G+1)^{-\nu} + (G+2)^{-\nu} + \dots]} > \frac{1}{1 + \gamma} > 1 - \gamma.$$

Auf Grund der vorstehenden Ungleichungen können wir jetzt in allen Fällen, wo $n > 4$ ist, die Beziehung $M_0 = 1$ nachweisen.

Wir wählen einfach die Zahl R derart, daß in ΔR sämtliche ungeraden Primzahlen auftreten, die kleiner als eine Zahl $G+1$ sind. Unsere Ungleichungen liefern uns dann, wenn n ungerade und > 3 ist, für alle Größen:

$$D_{\frac{n-1}{2}} \left[(-1)^{\frac{n-1}{2}} \sigma_1 m \Delta R^2 \right], \quad (8 \Delta R)_n \cdot S_n, \quad \frac{1}{(8 \Delta R)_{n-1} \cdot S_{n-1}},$$

und wenn n gerade und > 4 ist, für alle Größen:

$$\prod \left[1 + \left(\frac{(-1)^{\frac{n}{2}} \Delta R^2}{p} \right)^{\frac{1}{\frac{n-2}{2}}} \right], \quad (8 \Delta R)_n \cdot S_n, \quad \frac{1}{D_n \left[\frac{(-1)^{\frac{n}{2}} \Delta R^2}{2} \right]}$$

einmal obere und dann untere Grenzen. Indem wir erst diese oberen und dann diese unteren Grenzen einsetzen und jedesmal die hervorgehende Ungleichung durch die Gleichung (22) dividieren, bekommen wir, wenn $n \equiv 1 \pmod{2}$ und > 3 :

$$\frac{1 - \gamma_{\frac{n-3}{2}}}{1 + \gamma_{\frac{n-2}{2}}} < M_0 < \left(1 + \gamma_{\frac{n-3}{2}} \right) (1 + \gamma_{n-1}),$$

und wenn $n \equiv 0 \pmod{2}$ und > 4 :

$$\frac{1 - \gamma_{\frac{n-4}{2}}}{1 + \gamma_{\frac{n-2}{2}}} < M_0 < \frac{\left(1 + \gamma_{\frac{n-4}{2}} \right) (1 + \gamma_{n-1})}{1 - \gamma_{\frac{n-2}{2}}},$$

wobei γ_h für $\frac{1}{h G^h}$ gesetzt ist. Lassen wir jetzt die Zahl G ins Unendliche wachsen, so folgt in der Tat: $M_0 = 1$.

Es bleibt uns noch übrig, die Fälle $n = 2, 3, 4$ zu untersuchen.

Ist $n = 2$, so nehme man $R = 1$ und betrachte zu gleicher Zeit alle die Grenzwerte $L(m)$, welche die rechte Seite der Gleichung (19) darstellt, wenn man den $2 + \mathfrak{b}$ Einheiten $C(m)$ alle die Wertsysteme beilegt, die der Bedingung $\left(\frac{-\Delta}{m} \right) = 1$ genügen. Man bilde aus diesen $2^{1+\mathfrak{b}}$ Grenzwerten ebensoviele lineare Kombinationen: $\sum c(m) L(m) = L_c$; $c(m)$ bedeute hier ein beliebiges Glied des über alle Einheiten $C(m)$ erstreckten Produktes $\prod [1 + C(m)]$; von je zwei Einheiten $c(m)$, die ein Produkt gleich $\left(\frac{-\Delta}{m} \right)$ geben, soll aber immer nur eine beibehalten werden. Aus den Summen L_c folgen umgekehrt die Größen $L(m)$ mit Hilfe der Gleichungen $2^{1+\mathfrak{b}} \cdot L(m) = \sum_c c(m) \cdot L_c$. Die Grenzwerte L_c sind von Dirichlet angegeben. Unter ihnen ist nur einer von Null verschieden, nämlich derjenige, welcher zu $c = 1 = \left(\frac{-\Delta}{m} \right)$ gehört; dieser besitzt einen Ausdruck, aus welchem sofort $M_0 = 1$ hervorgeht.

In den Fällen $n = 3$ und $n = 4$ gebe ich für die Relation $M_0 = 1$ einen Beweis, welcher sich auf die Sätze aus der Anmerkung zu 9. stützt. Übrigens ließe sich für diese Fälle noch dieselbe Methode verwerten, welche in den Fällen $n > 4$ angewandt wurde. Für $n = 3$ sehe man auch: Smith, *On the Orders and Genera of Ternary Quadratic Forms*, artt. 13

—21 (Phil. Trans. CLVII, 1867; Collected Papers, vol. I) und: *F. Q.*, pp. 156—159 [[S. 127—129]]; für $n = 4$: *F. Q.*, pp. 162—163 [[S. 131—133]], und: Smith, *Sur la représentation des nombres par une somme de cinq carrés* (Mém. prés. T. XXIX; Collected Papers, vol. II).

Ist $n = 3$, so konstatiere man zunächst, daß dem speziellen Genus G von der Ordnung

$$\begin{pmatrix} 1, & 1 \\ 1, & 1 \end{pmatrix}, \quad (\Delta = 1)$$

welches die Formen von der Determinante 1 enthält, ein $M_0 = 1$, d. i. ein $M = \frac{1}{24}$ zukommt. Bekanntlich bilden alle Formen von G eine einzige Klasse, welche durch $f = x_1^2 + x_2^2 + x_3^2$ repräsentiert wird*), und dieses f läßt in der Tat genau 24 Transformationen von der Determinante 1 in sich selbst zu. Die Anwendung der Gleichung (21) auf G liefert:

$$(23) \quad \frac{(2R)_1}{2^{2+a}} \cdot \frac{(2R)_2 \cdot S_2}{(2R)_3 \cdot S_3} = \text{Lim} \left\{ \varrho \sum \frac{1}{m^{1+\varrho}} D_1[-mR^2] \right\},$$

wo m alle positiven und zu $2R$ relativ primen Zahlen mit festen Einheiten:

$$(-1)^{\frac{m-1}{2}}, \quad \left(\frac{2}{m}\right), \quad \left(\frac{m}{r_1}\right), \quad \left(\frac{m}{r_2}\right), \quad \dots, \quad \left(\frac{m}{r_r}\right)$$

zu durchlaufen hat. Dabei ist nach 9. Anm. die Einheit $(-1)^{\frac{m-1}{2}}$ stets gleich $+1$ zu nehmen, während die übrigen Einheiten ganz nach Belieben gewählt werden dürfen.

Die vorstehende Formel benutzen wir, um für ein beliebiges ternäres Genus G von einer Ordnung

$$\begin{pmatrix} \sigma_1, & \sigma_2 \\ o_1, & o_2 \end{pmatrix} \quad (\Delta = o_1^2 o_2)$$

die Relation $M_0 = 1$ abzuleiten. Nach (21) genügt die Größe M_0 eines solchen Genus jedenfalls einer Gleichung

$$(24) \quad \frac{(2\Delta)_1}{2^{2+b}} \cdot \frac{(2\Delta)_2 \cdot S_2}{(2\Delta)_3 \cdot S_3} \cdot M_0 = \text{Lim} \left\{ \varrho \sum \frac{1}{m^{1+\varrho}} D_1[-\sigma_1 \Delta m] \right\} = \{m\},$$

wo m alle positiven und zu 2Δ relativ primen Zahlen mit bestimmten festen Einheiten

$$C(m) \quad (-1)^{\frac{m-1}{2}}, \quad \left(\frac{2}{m}\right), \quad \left(\frac{m}{\partial_1}\right), \quad \left(\frac{m}{\partial_2}\right), \quad \dots, \quad \left(\frac{m}{\partial_b}\right)$$

durchläuft.

Wir betrachten zuerst den Fall, wo $\sigma_1 o_2$ und $\sigma_2 o_1$ beide vollständige Quadrate sind.

Das Genus G werde durch eine charakteristische Form f repräsen-

*) Vgl. z. B. Dirichlet in Crelles Journal, Bd. 40, S. 228. (Werke, Bd. II, S. 92.)

tiert. Der erste Koeffizient von f heie $\sigma_1 \varphi_1$, und es sei $\sigma_2 \varphi_2$ der erste Koeffizient der zu f adjungierten Form. Die Zahlen φ_1 und φ_2 sind dann prim zueinander, und es gelten zwei Kongruenzen

$$\begin{aligned} -o_1 \sigma_2 \varphi_2 &\equiv X_1^2 \pmod{\sigma_1^2 \varphi_1} \\ -o_2 \sigma_1 \varphi_1 &\equiv X_2^2 \pmod{\sigma_2^2 \varphi_2}. \end{aligned}$$

Dieselben geben

$$-\left(\frac{-\varphi_2}{\varphi_1}\right) \cdot \left(\frac{-\varphi_1}{\varphi_2}\right) = (-1)^{\frac{\varphi_1+1}{2} \cdot \frac{\varphi_2+1}{2}} = -1,$$

also

$$\varphi_1 \equiv 1, \quad \varphi_2 \equiv 1 \pmod{4},$$

was nur mit $\sigma_1 = 1, \sigma_2 = 1$ vertrglich ist. Denn htte man etwa $\sigma_2 = 2$, so wrde die zweite Kongruenz zugleich $-\varphi_1 \equiv 1 \pmod{4}$ liefern. Nach 7. besitzt nun unser Genus den Hauptrest $\varphi_1 \xi_1^2 + \frac{o_1 \varphi_2}{\varphi_1} \xi_2^2 + \frac{o_1 o_2}{\varphi_2} \xi_3^2 \pmod{4}$.

Derselbe zeigt, da in (24) die Einheit $(-1)^{\frac{m-1}{2}}$ allein gleich $+1$ genommen werden darf. Setzt man $\sigma_1 \Delta = 2^{2h} \cdot R^2$ ($R \equiv 1, \pmod{2}$), so kann daher der Grenzwert $\{m\}$ nach der Formel (23) bestimmt werden, und man findet $M_0 = 1$.

Zweitens sei eine der Zahlen $\sigma_1 o_2$ und $\sigma_2 o_1$ kein vollstndiges Quadrat.

Das Ma des Genus G stimmt, wie wir wissen, mit dem Mae des adjungierten Genus von der Ordnung

$$\begin{pmatrix} \sigma_2, \sigma_1 \\ o_2, o_1 \end{pmatrix}$$

berein; ebenso liefern diese beiden Genera gleiche Gren $f\{q\}$; sie mssen also auch dasselbe M_0 ergeben. Wir brauchen daher nur eines dieser Genera zu untersuchen und knnen annehmen, es sei $\sigma_1 o_2$, also auch $\sigma_1 \Delta$ kein vollstndiges Quadrat.

Aus $\sigma_1 \Delta$ gehe durch Division mit einer mglichst hohen Potenz von 4 die Zahl d hervor. Wir betrachten irgendein Genus φ_1 der Ordnung

$$\begin{pmatrix} 1, 1 \\ 1, d \end{pmatrix},$$

denken uns aber im Falle $d \equiv 1 \pmod{4}$ (was dann sicher gestattet ist), die Charaktere $\left(\frac{\varphi_2}{d}\right)$ dieses Genus so ausgesucht, da die Einheit $\delta = (-1)^{\frac{d+1}{2}} \cdot \left(\frac{\varphi_2}{d}\right) = (-1)^{\frac{\varphi_1 d+1}{2} \cdot \frac{\varphi_2+1}{2}}$ gleich $+1$ wird. Die zu φ_1 gehrige Gre M_0 lt sich ebenfalls durch einen der Grenzwerte $\{m\}$ darstellen; und zwar ist hier, nach den Stzen aus 9. Anm., ein jeder der 2^{2+b} Grenzwerte $\{m\}$ in gleicher Weise verwendbar. Alle die Grenzwerte $\{m\}$ mssen demnach untereinander gleich sein.

Bilden wir jetzt die Formel (24) für das zu φ_1 adjungierte Genus φ_2 der Ordnung

$$\begin{pmatrix} 1, & 1 \\ d, & 1 \end{pmatrix},$$

so ergeben sich die Grenzwerte $\{m\}$ gleich gewissen Grenzwerten

$$\text{Lim} \left\{ \varrho \sum \frac{1}{m^{1+\varrho}} D_1[-d^2 m] \right\},$$

wo m wie vorher Reihen von positiven Zahlen mit festen Charakteren $C(m) = \pm 1$ zu durchlaufen hat. Hier ist für die Einheit $(-1)^{\frac{m-1}{2}}$, nach 9. Anm., stets der Wert $+1$ zulässig. Wenn man $d = R$, resp. $= 2R$ setzt, kann man also für den vorstehenden Grenzwert die Formel (23) benutzen, und man gelangt zu $M_0 = 1$.

Ist endlich $n = 4$, so haben wir einen Grenzwert

$$\text{Lim} \left[\varrho \sum \frac{1}{m^{1+\varrho}} \prod_p \left\{ 1 + \left(\frac{-\Delta}{p} \right) \frac{1}{p} \right\} \right] = \{m\}$$

zu ermitteln, wo m alle positiven und zu 2Δ relativ primen Zahlen mit festen Einheiten

$$C(m) \quad (-1)^{\frac{m-1}{2}}, \quad \left(\frac{2}{m} \right), \quad \left(\frac{m}{\partial_1} \right), \quad \left(\frac{m}{\partial_2} \right), \quad \dots, \quad \left(\frac{m}{\partial_b} \right)$$

durchläuft. Wir bilden aus Δ durch Division mit einer möglichst hohen Potenz von 4 eine Zahl Δ_0 . Die Größe M_0 für irgendein Genus der Ordnung

$$\begin{pmatrix} 1, & 1, & 1 \\ 1, & 1, & \Delta_0 \end{pmatrix}$$

hängt dann gleichfalls von irgendeinem Grenzwerte $\{m\}$ ab. Für ein solches Genus unterliegen aber, nach den Sätzen aus 9. Anm., die Einheiten $C(m)$ durchaus keiner Beschränkung. Alle die 2^{2+b} Grenzwerte $\{m\}$ müssen also untereinander gleich sein, und wir können sie gleich dem 2^{2+b} ten Teile ihrer Summe setzen. Diese Summe wird durch einen ähnlichen Grenzwert gebildet, wo m alle möglichen positiven und zu 2Δ relativ primen Zahlen durchläuft. Dieser letzte Grenzwert gestattet nach F. Q., p. 150 und 162 [[S. 123 und 132]] eine Transformation in:

$$\text{Lim} \left(\varrho \sum \frac{1}{m^{1+\varrho}} \cdot \frac{\sum \left(\frac{\Delta}{m} \right) m^{-2-\varrho}}{\sum m^{-4-2\varrho}} \right) = (2\Delta)_1 \cdot \frac{D_2[\Delta]}{(2\Delta)_4 \cdot S_4},$$

welche dann unmittelbar zu $M_0 = 1$ führt.

Von den verschiedenen Darstellungen, deren das Maß eines Genus fähig ist, erscheint als die natürlichste die hier gegebene mit Hilfe eines unendlichen Produktes, in welchem einer jeden Primzahl ein bestimmter

Faktor entspricht*). Ihre Bedeutung reicht über das spezielle Gebiet der quadratischen Formen hinaus: es zeigt diese Darstellung, daß zur Lösung arithmetischer Probleme über Formenanzahlen ein Studium jener wichtigen Gruppenbildungen erforderlich ist, auf welche Herr Camille Jordan in Nr. 302 des *Traité des Substitutions* aufmerksam gemacht hat.

Ausdrücke für das Maß eines positiven Genus quadratischer Formen von n Variablen sind zuerst von Henry J. Stephen Smith in der Note *On the Orders and Genera of Quadratic Forms containing more than three Indeterminates* (Roy. Soc. Proc. XVI, 1868, pp. 197—208; Collected Papers, vol. I, pp. 510—523) mitgeteilt. Die Formeln von Smith sind ähnlich unseren Formeln (17) in 8., erschöpfen aber nicht alle Fälle; sie geben im wesentlichen die Werte der von uns E_q genannten Faktoren für ungerade Primzahlen q , doch für die Primzahl 2 nur in den weniger verwickelten Fällen, wo das Genus eine ungerade Determinante besitzt.

In einem folgenden Aufsätze beabsichtige ich verschiedene Anwendungen der hier gefundenen Resultate auseinanderzusetzen.

Vita.

Natus sum Hermann Minkowski in Russiae oppido Alexoten die XXII. mensis junii anni 1864. Anno 1872 in civitatem Borussorum susceptus, gymnasium Palaeopolitanum Regimonti adii. Vere anni 1880 maturitatis testimonium adeptus, per quinque semestria Regimonti, per tria Berolini studiis mathematicis incubui. Docuerunt me viri illustrissimi: de Helmholtz, Hurwitz, Lindemann, Kirchhoff, Kronecker, Kummer, Runge, Voigt, Weber, Weierstraß, quibus omnibus optime de me meritis gratias ago maximas.

Thesen.

I. Es ist nicht wahrscheinlich, daß eine jede positive Form sich als eine Summe von Formenquadraten darstellen läßt.

II. Es hat seine Bedenken, in mathematischen Untersuchungen sich auf räumliche Anschauung zu berufen.

Öffentliche Verteidigung der Dissertation und der Thesen: am 30. Juli 1885.

Opponenten: Herr Dr. David Hilbert,
Herr Emil Wiechert, stud. math.

*) Ich habe auf diese Darstellung im 99. Bande des Journals für die reine und angewandte Mathematik hingewiesen. Diese Ges. Abhandlungen S. 150—153.

V.

Über den arithmetischen Begriff der Äquivalenz und über die endlichen Gruppen linearer ganzzahliger Substitutionen.

(Crelles Journal für die reine und angewandte Mathematik, Band 100, S. 449—458).

Herr Kronecker stellt in § 24 (S. 107) seiner *Festschrift zu Herrn Kummers Doktor-Jubiläum**) an eine rationelle Fassung des arithmetischen Begriffs der Äquivalenz von Formen die Forderungen, daß auf Grund einer solchen die verschiedenen Formenklassen gleich dicht werden, und die Anzahl der Klassen möglichst klein ausfalle, und entwickelt in §§ 1 und 2 der Abhandlung: *Über bilineare Formen mit vier Variabeln***), wie diesen Forderungen für definite quadratische Formen mit zwei Variablen genügt werden kann. Dabei erscheint der Umstand von wesentlicher Bedeutung, daß diese Formen, abgesehen von der identischen und der negativen identischen Transformation, keine eigentlichen Transformationen in sich zulassen können, welche modulo 2 der identischen Transformation kongruent sind.

Im folgenden will ich zeigen, in welcher Weise die Kroneckerschen Forderungen für alle diejenigen homogenen ganzen Formen irgendeiner Variablenzahl zu erfüllen sind, welche nur eine endliche Zahl von linearen ganzzahligen Transformationen in sich zulassen, also beispielsweise für alle wesentlich positiven (oder wesentlich negativen) quadratischen Formen von nichtverschwindender Determinante.

§ 1.

Es gibt keine homogene lineare ganze ganzzahlige Substitution mit n Variablen von einer endlichen Ordnung, welche modulo 4 der identischen Substitution kongruent wäre, diese selbst ausgenommen.

Denn es sei

$$A: x_h = a_{h1}y_1 + a_{h2}y_2 + \cdots + a_{hn}y_n \quad (h = 1, 2, \dots, n)$$

*) Crelles Journal, Bd. 92. (Werke, Bd. II, S. 370.)

**) Abhandlungen der Berliner Akademie, 1883. (Werke, Bd. II, S. 429.)

eine solche Substitution, es mögen also die Kongruenzen gelten:

$$a_{hh} \equiv 1, \quad a_{hk} \equiv 0 \pmod{4} \quad (h \neq k).$$

Damit A eine endliche Ordnung besitze, ist notwendig und hinreichend, daß die Elementarteiler der mit einem Parameter r gebildeten charakteristischen Determinante

$$\Delta = (-1)^n \cdot |a_{hk} - r\delta_{hk}| \quad \left(\begin{array}{l} h, k = 1, 2, \dots, n, \\ \delta_{hh} = 1, \delta_{hk} = 0, h \neq k \end{array} \right)$$

nur für Wurzeln der Einheit verschwinden und sämtlich linear seien*).

Bedeutet $f_\nu(r)$, für irgendeine ganze Zahl $\nu > 1$, diejenige irreduzible ganze Funktion $\varphi(\nu)^{\text{ten}}$ Grades, welche für die primitiven ν^{ten} Einheitswurzeln Null wird, und deren höchster Koeffizient gleich 1 ist, so wird also die Determinante Δ jedenfalls einen Ausdruck erhalten

$$(1) \quad (r-1)^m f_{\nu_1}(r) f_{\nu_2}(r) \dots f_{\nu_g}(r), \quad (m \geq 0, \nu > 1)$$

wo die ganzen Zahlen m, ν der Bedingung genügen werden

$$(2) \quad n = m + \varphi(\nu_1) + \varphi(\nu_2) + \dots + \varphi(\nu_g) \geq m + g.$$

Man setze nun für r irgendeine ganze Zahl

$$c \equiv 1 + 4 \pmod{8}$$

ein. Dann muß die Determinante Δ durch 2^{2n} aufgehen, weil jeder ihrer Koeffizienten durch 4 teilbar wird.

Untersuchen wir die höchste Potenz von 2, welche in (1) erscheint. Geht in einer Zahl ν die Potenz 2^x , aber nicht mehr 2^{x+1} auf, so folgt

$$c^\nu \equiv c^{2^x} \equiv 1 + 2^{x+2} \pmod{2^{x+3}};$$

also wird $c^\nu - 1$ genau die Potenz 2^{x+2} aus 4ν enthalten. Nun hat ein $f_\nu(r)$ bekanntlich den Ausdruck

$$\frac{(r^\nu - 1)(r^{\frac{\nu}{\alpha}} - 1)(r^{\frac{\nu}{\beta}} - 1) \dots}{(r^\alpha - 1)(r^\beta - 1)(r^\gamma - 1) \dots},$$

wenn $\alpha, \beta, \gamma, \dots$ die verschiedenen Primzahlen aus ν vorstellen; also enthält $f_\nu(c)$ dieselbe Potenz von 2 wie

$$\frac{4\nu \cdot 4 \cdot \frac{\nu}{\alpha} \cdot 4 \cdot \frac{\nu}{\beta} \cdot 4 \cdot \frac{\nu}{\gamma} \cdot \dots}{4 \cdot \frac{\nu}{\alpha} \cdot 4 \cdot \frac{\nu}{\beta} \cdot 4 \cdot \frac{\nu}{\gamma} \cdot \dots}.$$

Die letztere Zahl ist gleich 1, wenn die Zahl ν sich aus mehreren verschiedenen Primzahlen zusammensetzt, dagegen, wenn ν Potenz einer einzigen Primzahl ist, gleich dieser Primzahl, also entweder ungerade oder

*) Hermite, Crelles Journal, Bd. 47, S. 312 (Oeuvres, T. I, p. 199); C. Jordan, Crelles Journal, Bd. 84, S. 112.

gleich 2. Irgendein $f_v(c)$ ist also höchstens durch 2, das ganze Produkt (1) für $r=c$ also höchstens durch die Potenz 2^{2m+g} teilbar, die wegen (2) stets $< 2^{2n}$ ausfällt, außer im Falle $m=n$, wo dann überhaupt kein Faktor $f_v(r)$ in Δ auftritt.

Demnach verschwinden die Elementarteiler von Δ nur für $r=1$, und da sie sämtlich linear sein sollen, so müssen in Δ für $r=1$ selbst die Koeffizienten verschwinden, d. h. A ist die identische Substitution, w. z. b. w.

Ein Quotient aus zwei verschiedenen, modulo 4 kongruenten ganzzahligen Substitutionen von endlicher Ordnung ist eine ganzzahlige Substitution von unendlicher Ordnung.

§ 2.

Zwei ganzzahlige Substitutionen A und $S^{-1}AS$ sollen im folgenden nur dann als *ähnlich* bezeichnet werden, wenn die transformierende Substitution S ganzzahlig und von einer Determinante ± 1 ist.

Ist A modulo 2 der identischen Substitution kongruent, so gilt dasselbe von jeder ähnlichen Substitution.

Hat für eine Substitution A die charakteristische Determinante Δ einen Ausdruck

$$(r-1)^m(r+1)^{n-m},$$

dann gibt es ähnliche Substitutionen, in welchen:

$$a_{hk}=0 \quad (h < k), \quad a_{hh}=1 \quad (h \leq m), \quad a_{hh}=-1 \quad (h > m).$$

Denn ist A nicht selbst eine solche Substitution, so sei in A ν der kleinste Wert von h , für welchen diese Gleichungen nicht mehr statthaben. Dann bildet die Determinante

$$r\delta_{hk} - a_{hk} \quad (h, k = \nu, \dots, n)$$

einen Divisor von Δ und verschwindet, wenn $\nu \leq m$, noch für $r = \varepsilon = 1$, wenn $\nu > m$, für $r = \varepsilon = -1$. Man kann daher $n - \nu + 1$ ganze Zahlen $s_\nu, \dots, s_n$ ohne gemeinsamen Teiler bestimmen, so daß

$$\varepsilon s_k = \sum_h a_{hk} s_h \quad (h, k = \nu, \dots, n)$$

wird, und ferner eine Substitution S mit einer Determinante ± 1 , so daß man in S^{-1}

$$x_h = y_h \quad (h < \nu), \quad x_\nu = \sum_{h=\nu}^n s_h y_h$$

hat.*) Dann finden sich in der Substitution $S^{-1}AS$ die obigen Gleichungen auch für $h = \nu$ erfüllt, und es ist evident, wie man weiter zu schließen hat.

*) Diese Stelle ist nach einer von Minkowski selbst herrührenden Berichtigung (Crelles Journal, Bd. 101, S. 202) korrigiert. (Anm. des Herausg.)

§ 3.

Es gibt, außer den Substitutionen, welche einer Substitution

$$x_h = \varepsilon_h y_h; \quad \varepsilon_h = +1, -1 \quad (h = 1, 2, \dots, n)$$

ähnlich sind, keine ganzzahlige Substitution von einer endlichen Ordnung, welche modulo 2 der identischen Substitution kongruent wäre.

Denn es bedeute A eine solche Substitution, es sei also

$$a_{hh} \equiv 1, \quad a_{hk} \equiv 0 \pmod{2} \quad (h \neq k).$$

Dann ist $A \cdot A$, wie man sich leicht überzeugt, der identischen Substitution modulo 4 kongruent, muß also, nach § 1, mit dieser übereinstimmen. Die Elementarteiler der charakteristischen Determinante von A können daher nur für $r = 1$ oder $r = -1$ verschwinden, und diese Determinante hat einen Ausdruck

$$(r - 1)^m (r + 1)^{n-m}.$$

Nach § 2 kann man daher eine Substitution A^* bestimmen, welche A ähnlich ist und in welcher

$$a_{hk}^* = 0 \quad (h < k), \quad a_{hh}^* = 1 \quad (h \leq m), \quad a_{hh}^* = -1 \quad (h > m)$$

ist. Damit die Elementarteiler von A^* linear ausfallen, muß noch

$$a_{hk}^* = 0 \quad (m \leq h < k), \quad a_{hk}^* = 0 \quad (h > k > m)$$

sein, so daß man in A^* hat:

$$x_h = y_h \quad (h \leq m), \quad x_h = \sum 2p_{hk} y_k - y_h \quad (h > m, \quad k = 1, \dots, m).$$

Dann findet man $A^* = P^{-1} \mathfrak{A} P$, wenn P die Substitution

$$x_h = y_h \quad (h \leq m), \quad x_h = -\sum p_{hk} y_k + y_h \quad (h > m, \quad k = 1, \dots, m)$$

und $\mathfrak{A}$ die Substitution

$$x_h = y_h \quad (h = 1, \dots, m), \quad x_h = -y_h \quad (h = m + 1, \dots, n)$$

bedeutet, womit unser Satz bewiesen ist.

Ist $0 < m < n$, so zerlegt sich jede quadratische Form, welche durch $\mathfrak{A}$ in sich selbst transformiert wird, in $f_1 + f_2$, wo f_1 nur von $x_1, \dots, x_m$ und f_2 nur von $x_{m+1}, \dots, x_n$ abhängt.

Um zu erkennen, inwieweit die gefundene Darstellung von A willkürlich bleibt, untersuchen wir eine Gleichung:

$$S^{-1} \mathfrak{A} S = \mathfrak{A}, \quad S = \pm 1.$$

Man denke sich S und $\mathfrak{A}$ als bilineare Formen mit zwei Reihen von n Variablen, teile jede Reihe in m und $n - m$ Variablen ein, und zerlege entsprechend S in vier Teile und $\mathfrak{A}$ in $\mathfrak{A}_1 - \mathfrak{A}_2$, dann ergibt sich

$$(S_{11} + S_{12} + S_{21} + S_{22})(\mathfrak{A}_1 - \mathfrak{A}_2) = (\mathfrak{A}_1 - \mathfrak{A}_2)(S_{11} + S_{12} + S_{21} + S_{22}),$$

$$S_{11} - S_{12} + S_{21} - S_{22} = S_{11} + S_{12} - S_{21} - S_{22},$$

also

$$S_{12} = 0, \quad S_{21} = 0 \quad \text{und} \quad S = S_{11} + S_{22}, \quad S_{11} = \pm 1, \quad S_{22} = \pm 1.$$

§ 4.

Es bedeute wieder A eine, modulo 2 der identischen kongruente Substitution, und es werde gesetzt: $a_{hh} \equiv \varepsilon_h \pmod{4}$, $\varepsilon_h = \pm 1$. Dann sind in εA , wenn ε die Substitution

$$x_h = \varepsilon_h y_h \quad (h = 1, 2, \dots, n)$$

vorstellt, alle ungeraden Koeffizienten $\equiv 1 \pmod{4}$. Hat ferner letzteres für irgend zwei Substitutionen statt, welche modulo 2 der identischen Substitution kongruent sind, so findet es sich auch im Produkte derselben erfüllt.

Jede ganzzahlige Substitution

$$A: \quad x_h = \sum a_{hk} y_k \quad (h, k = 1, 2, \dots, n)$$

von einer Determinante 1, welche den Kongruenzen

$$a_{hh} \equiv 1 \pmod{4}, \quad a_{hk} \equiv 0 \pmod{2} \quad (h \neq k)$$

genügt, läßt sich als Produkt von lauter Substitutionen $S_{hk}^{\pm 1}$ von der Form:

$$x_1 = y_1, \dots, x_{h-1} = y_{h-1}, \quad x_h = y_h \pm 2y_k, \quad x_{h+1} = y_{h+1}, \dots, x_n = y_n \quad (h \neq k)$$

darstellen.

Man kann beim Beweise dieses Satzes sich einer Methode bedienen, ähnlich derjenigen, welche Herr Kronecker für die Zerlegung beliebiger ganzzahliger Substitutionen in elementare aufgestellt hat*).

Es genügt, zu zeigen, daß A durch wiederholtes Zusammensetzen mit Substitutionen $S_{hk}^{\pm 1}$ auf die identische Substitution reduziert werden kann. Sei in A x_h die erste Variable, welche nicht einfach in y_h übergehe, dann ergibt die Determinante von A :

$$1 \equiv 0 \pmod{a_{hh}, a_{h,h+1}, \dots, a_{hn}},$$

also $a_{hh} = 1$, sobald alle Zahlen $a_{h,h+1}, \dots, a_{hn}$ verschwinden.

Sei aber noch eine dieser Zahlen, z. B. a_{hk} ($h < k$), von Null verschieden, und bedeute ± 1 die positive oder negative Einheit, je nachdem a_{hh} und a_{hk} gleiches oder verschiedenes Zeichen haben. Der absolute Wert der ungeraden Zahl a_{hh} ist verschieden von dem absoluten Werte der geraden Zahl a_{hk} ; ist er kleiner als letzterer, so finden sich die Zahlen a_{hh} und a_{hk} in $AS_{hk}^{\pm 1}$ durch a_{hh} , $a_{hk} \mp 2a_{hh}$, ist er größer, in $AS_{hk}^{\mp 1}$ durch $a_{hh} \mp 2a_{hk}$, a_{hk} , also jedesmal durch zwei Zahlen ersetzt, von denen eine unverändert, die andere dem absoluten Werte nach kleiner ist. So können a_{hh} , $a_{h,h+1}, \dots, a_{hn}$ allmählich auf 1, 0, $\dots$, 0 reduziert werden.

Ist alsdann eine der Zahlen $a_{h1}, \dots, a_{h,h-1}$, etwa a_{hk} ($h > k$), von Null verschieden und ± 1 ihr Vorzeichen, so erscheint diese Zahl in $AS_{hk}^{\mp 1}$

*) Monatsberichte der Berliner Akademie, 1866, S. 608 ff., oder Crelles Journal, Bd. 68, S. 282 ff. (Werke, Bd. I, S. 158 ff.)

durch $a_{hk} \mp 2$ ersetzt, also, absolut genommen, verkleinert. Alle diese Zahlen können daher zum Verschwinden gebracht werden.

In der auf solche Weise reduzierten Substitution erfährt auch die h^{te} Variable keine Änderung, und damit ist, unter Zuhilfenahme eines Schlusses von h auf $h + 1$, unser Satz verifiziert.

§ 5.

1. Es sei G irgendeine endliche Gruppe von homogenen linearen ganzen ganzzahligen Substitutionen mit n Variablen. In derselben bilden diejenigen Substitutionen, welche modulo 2 der identischen Substitution kongruent sind, eine ausgezeichnete Untergruppe, welche $\mathfrak{G}$ heißen möge.

Die Ordnung von $\mathfrak{G}$ ist eine Potenz 2^n , $0 \leq n \leq n$. Die Gruppe $\mathfrak{G}$ ist ähnlich einer Untergruppe der Gruppe der 2^n Substitutionen

$$\mathfrak{S}: \quad x_h = \pm y_h. \quad (h = 1, 2, \dots, n)$$

Eine jede Substitution $\mathfrak{S}$ von $\mathfrak{G}$ genügt nämlich nach § 3 der Gleichung $\mathfrak{S}\mathfrak{S} = \mathfrak{E}$, wenn $\mathfrak{E}$ die identische Substitution bedeutet.

Findet man nun in $\mathfrak{G}$ außer $\mathfrak{E}$ (und eventuell $-\mathfrak{E}$) noch eine Substitution $\mathfrak{A}$, so kann man, nach § 3, $\mathfrak{G}$ derart transformiert voraussetzen, daß $\mathfrak{A}$ einen Ausdruck hat

$$x_h = y_h \quad (h = 1, \dots, m), \quad x_h = -y_h \quad (h = m + 1, \dots, n).$$

Erscheint dann in $\mathfrak{G}$ außer der Gruppe $\{\mathfrak{E}, -\mathfrak{E}, \mathfrak{A}\}$ noch ein $\mathfrak{B}$, so ist auch $\mathfrak{B}\mathfrak{A} \cdot \mathfrak{B}\mathfrak{A} = \mathfrak{E}$, also $\mathfrak{B}^{-1}\mathfrak{A}\mathfrak{B} = \mathfrak{A}$. Nach § 3 muß daher $\mathfrak{B}$ sich in $\mathfrak{B}_1 + \mathfrak{B}_2$ zerlegen, wo $\mathfrak{B}_1$ und $\mathfrak{B}_2$ Substitutionen mit m und $n - m$ Variablen sind. Dieselben sind ebenso wie $\mathfrak{B}$ von einer endlichen Ordnung und modulo 2 von der Ordnung Eins, können daher, nach § 3, durch Transformation mit Substitutionen S_1 und S_2 von einer Determinante ± 1 in einen, $\mathfrak{J}$ analogen Typus übergeführt werden. Durch Transformation mit $S_1 + S_2$ erlangt dann auch $\mathfrak{B}$ den Typus $\mathfrak{J}$, während $\mathfrak{A}$ sowie $\mathfrak{E}$ ungeändert bleiben.

Enthält jetzt $\mathfrak{G}$ außer der Gruppe $\{\mathfrak{E}, -\mathfrak{E}, \mathfrak{A}, \mathfrak{B}\}$ noch ein $\mathfrak{C}$, so ist $\mathfrak{C}^{-1}\mathfrak{A}\mathfrak{C} = \mathfrak{A}$, $\mathfrak{C}^{-1}\mathfrak{B}\mathfrak{C} = \mathfrak{B}$ usw. Sollte endlich $\mathfrak{G}$, auf irgendeine Weise transformiert, alle Substitutionen $\mathfrak{J}$ enthalten, so würden damit sicher die Substitutionen von $\mathfrak{G}$ erschöpft sein. Denn ein jedes $\mathfrak{S}$ müßte dann, wegen $\mathfrak{S}^{-1}\mathfrak{J}\mathfrak{S} = \mathfrak{J}$, sich auf alle möglichen Arten in $\mathfrak{S}_1 + \mathfrak{S}_2$ zerlegen lassen.

Die Substitutionen von G verteilen sich in Reihen $\mathfrak{G}, A\mathfrak{G}, \dots$ von je 2^n Substitutionen, welche untereinander modulo 2 kongruent sind.

Die Gruppe G ist 2^n -stufig isomorph zur Gruppe G_R der Reste ihrer Substitutionen nach dem Modul 2. Letztere Gruppe bildet offenbar eine Untergruppe der Gruppe sämtlicher inkongruenter Substitutionenreste

modulo 2 von einer Determinante $\equiv 1 \pmod{2}$; ihre Ordnung ist daher ein Divisor der Ordnung dieser Gruppe, d. i. der Zahl:

$$N = (2^n - 1)(2^n - 2) \cdots (2^n - 2^{n-1}).$$

Also:

Die Ordnung jeder endlichen Gruppe von homogenen linearen ganzen ganzzahligen Substitutionen mit n Variablen ist ein Divisor der Zahl $2^n N$.

Ein eingehenderes Studium der endlichen Gruppen ganzzahliger Substitutionen, welches ich mir für einen folgenden Aufsatz vorbehalte, führt zu fundamentalen Beziehungen zwischen der Theorie der aus Wurzeln der Einheit gebildeten komplexen Zahlen und der Theorie der Reduktion der wesentlich positiven quadratischen Formen.

2. Sei jetzt f eine homogene ganze Form mit n Variablen, welche nur eine endliche Anzahl von linearen ganzzahligen Transformationen in sich zulasse; und sei $t(f)$ diese Anzahl.

Die $t(f)$ Transformationen werden eine endliche Gruppe bilden; also ist $t(f)$ ein Divisor von $2^n N$, also $t(f) \leq 2^n N$. Sie sind sämtlich von einer Determinante ± 1 ; man findet unter ihnen außer der identischen Transformation keine, welche der identischen modulo 4 kongruent wäre, dagegen noch irgend $2^n - 1$ Transformationen ($0 \leq n \leq n$), welche der identischen modulo 2 kongruent sind.

Die letzteren Transformationen können im Falle $n = 2$ nur sein: 1. wenn f von gerader Dimension ist, die negative identische Transformation, 2. Transformationen, welche der Transformation

$$x_1 = y_1, \quad x_2 = -y_2$$

ähnlich, also von einer Determinante -1 sind.

§ 6.

Es bedeute allgemein

S_a eine (homogene lineare ganze) ganzzahlige Substitution von einer Determinante ± 1 ,

S_e eine ebensolche Substitution von einer Determinante 1,

S_u eine ebensolche Substitution von einer Determinante 1, welche modulo 2 der identischen Substitution kongruent ist,

S_v , wenn $n = 2$ ist, eine Substitution $S_u = \begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix}$, in welcher entweder

$$\alpha - 1 \equiv \beta \equiv \gamma \equiv \delta - 1 \equiv 0 \pmod{4},$$

oder

$$\alpha + \delta \equiv -2 \pmod{8}$$

und dabei nicht

$$\alpha + 1 \equiv \beta \equiv \gamma \equiv \delta + 1 \equiv 0 \pmod{4}$$

ist,

wenn $n > 2$ ist, eine ganzzahlige Substitution von einer Determinante 1, welche der identischen Substitution modulo 4 kongruent ist.

Die S_a , die S_e , die S_u , alle S_v bilden Gruppen. Von solchen vier Gruppen ist jede folgende eine ausgezeichnete Untergruppe jeder vorhergehenden. Im Falle $n = 2$ entsteht die Gruppe der S_u aus der Gruppe der S_v durch Hinzunahme der Substitution

$$x_1 = -y_1, \quad x_2 = -y_2.$$

Zwei binäre Formen gerader Dimension, welche durch ein S_u ineinander übergehen, lassen sich daher auch durch ein S_v ineinander überführen.

Zwei homogene ganze Formen heißen *äquivalent*, wenn sie durch ein S_a , *eigentlich äquivalent*, wenn sie durch ein S_e ineinander transformiert werden können.

Zwei Formen sollen *vollständig**) *äquivalent* heißen, wenn sie durch ein S_v ineinander transformiert werden können.

Wir wollen uns auf Formen beschränken, welche nur eine endliche Anzahl von linearen ganzzahligen Transformationen in sich zulassen. Jede solche Form ist *sich selbst* nur vermittels der identischen Transformation vollständig äquivalent. Daher enthält eine Klasse vollständig äquivalenter Formen ebensoviel verschiedene Formen, als verschiedene Substitutionen S_v existieren. Der Begriff der vollständigen Äquivalenz führt also, der Kroneckerschen Forderung entsprechend, zu einer für alle Klassen konstanten, nur von der Variablenzahl n abhängenden Dichtigkeit.

Geht eine Form f durch $t(f)$ Transformationen in sich selbst über, so wird die Klasse der mit f äquivalenten Formen, je nachdem $n = 2$ oder > 2 ist, sich in $\frac{2^2 N}{t(f)}$ oder in $\frac{2^{n^2} N}{t(f)}$ verschiedene Klassen von untereinander vollständig äquivalenten Formen auflösen. N bedeutet hier die Zahl aus § 5. Da $t(f)$ stets in $2^n N$ aufgeht, so werden in den Fällen $n \geq 3$ die Klassenanzahlen den Faktor 2^{n^2-n} enthalten.

Auf Grund der Untersuchungen von Herrn Camille Jordan über die Zusammensetzung der linearen Gruppe (*Traité des Substitutions*, 127—140) läßt sich ferner nachweisen, daß man nicht durch eine von der hier gegebenen abweichende Fassung des Begriffs der vollständigen Äquivalenz zu kleineren Klassenanzahlen, d. i. zu einer größeren konstanten Dichtigkeit in den Klassen gelangen kann, wenn man nur folgendes voraussetzt: Die Gesamtheit der Substitutionen, welche vollständige Äquivalenz hervorrufen, soll eine *ausgezeichnete* Untergruppe der Gruppe derjenigen Substitutionen sein, welche Äquivalenz hervorrufen; oder (was dasselbe ist):

*) Kronecker, Über bilineare Formen mit vier Variablen, S. 9. (Werke, Bd. II, S. 434.)

Zwei vollständig äquivalente Formen sollen, irgendeiner Äquivalenz-erzeugenden Transformation zu gleicher Zeit unterworfen, vollständig äquivalent bleiben.

Wollte man indes von dieser Voraussetzung absehen, so könnte man in den Fällen $n \geq 3$ dem Begriffe der vollständigen Äquivalenz z. B. die Gruppe derjenigen Substitutionen von der Determinante 1 zugrunde legen, in welchen

$$a_{hh} \equiv 1 \pmod{4}, \quad a_{hk} \equiv 0 \pmod{4} \quad (h < k), \quad a_{hk} \equiv 0 \pmod{2} \quad (h > k)$$

ist; dann würden sich Klassen von einem $2^{\frac{n^2-n}{2}}$ -mal so großen Formen-inhalte ergeben.

Königsberg i. Pr., 1886.

VI.

Zur Theorie der positiven quadratischen Formen.*)

(Crelles Journal für die reine und angewandte Mathematik, Band 101, S. 196—202).

Eine wesentlich positive quadratische Form von n Variablen, mit reellen Koeffizienten und nichtverschwindender Determinante, kann — wie eine Darstellung der Form als Summe der Quadrate von n unabhängigen reellen linearen Formen leicht erkennen läßt — nur bei einer endlichen Anzahl ganzzahliger linearer Transformationen ungeändert bleiben. Jede einzelne von diesen Transformationen muß daher eine endliche Ordnung besitzen, d. h. nach einer endlichen Reihe von Wiederholungen zur identischen Transformation führen, und kann deshalb, nach § 1 meines Aufsatzes *Über den arithmetischen Begriff der Äquivalenz*, niemals der identischen Transformation modulo 4 kongruent sein, wenn sie nicht mit derselben übereinstimmt.

Das Gleiche gilt nun in bezug auf eine jede ungerade Primzahl als Modul; und aus diesem Umstande ergeben sich einige Aufschlüsse über die gesamte Anzahl der in Rede stehenden Transformationen, eine Anzahl, von welcher zuerst Herr Camille Jordan bewiesen hat, daß sie eine nur von der Zahl n abhängende Grenze nicht überschreiten kann.**)

§ 1.

Eine lineare Transformation

$$A: \quad x_h = a_{h1}y_1 + a_{h2}y_2 + \cdots + a_{hn}y_n \quad (h = 1, 2, \dots, n)$$

von endlicher Ordnung ist dadurch charakterisiert, daß die mit einem Parameter r gebildete Determinante

$$\Delta = |r\delta_{hk} - a_{hk}| \quad \left(\begin{array}{c} h, k = 1, 2, \dots, n \\ \delta_{hh} = 1, \delta_{hk} = 0, \quad h \neq k \end{array} \right)$$

nur für Einheitswurzeln verschwindet, und zwar für einen mehrfachen,

*) Der nachfolgende Aufsatz wurde in Verbindung mit dem Aufsätze *Über den arithmetischen Begriff der Äquivalenz* (Crelles Journal, Bd. 100, S. 449—458; diese Ges. Abhandlungen, Bd. I, S. 201—211) vom Verfasser im März 1887 der philosophischen Fakultät zu Bonn als Habilitationsschrift vorgelegt.

**) Journal de l'École polytechnique, cah. 48, p. 133.

etwa m -fachen Nullwert zusammen mit allen ihren $(m-1)^{\text{ten}}$ Unterdeterminanten*).

Bei ganzzahligen Koeffizienten a_{hk} liefert daher eine Zerlegung in irreduktible Faktoren für Δ einen Ausdruck:

$$(1) \quad (r-1)^m f_{\nu}(r) f_{\nu''}(r) \dots, \quad (m \geq 0, \nu > 1)$$

wenn mit $f_{\nu}(r)$ für eine ganze Zahl $\nu > 1$ diejenige ganze Funktion $\varphi(\nu)^{\text{ten}}$ Grades bezeichnet wird, welche für die primitiven ν^{ten} Einheitswurzeln verschwindet und als Koeffizient des höchsten Gliedes die Zahl 1 hat. Der Grad von (1) ist:

$$n = m + \varphi(\nu') + \varphi(\nu'') + \dots$$

Soll die Transformation A in bezug auf irgendeine ungerade Primzahl p der identischen Transformation kongruent sein, so muß sie mit derselben zusammenfallen.

Denn ist

$$a_{hk} \equiv \delta_{hk} \pmod{p}, \quad (h, k = 1, 2, \dots, n)$$

und setzt man für r eine ganze Zahl

$$c \equiv 1 + p \pmod{p^2},$$

so geht Δ durch p^n auf. Damit aber der Ausdruck (1) durch p^n teilbar werde, ist notwendig, daß in Δ kein $f_{\nu}(r)$ ($\nu > 1$) auftrete, daß also $\Delta = (r-1)^n$ sei. Denn $(c-1)^m$ enthält zwar genau p^m , irgendein $f_{\nu}(c)$ aber, wenn $\nu > 1$ ist, niemals die Potenz $p^{\varphi(\nu)}$.

Letzteres sieht man in folgender Weise ein. Ist eine Zahl ν ein Vielfaches von p^x , aber nicht mehr von p^{x+1} , so findet man:

$$c^{\nu} \equiv 1 + \nu p \pmod{p^{x+2}};$$

also enthält $c^{\nu} - 1$ dieselbe Potenz von p wie $p\nu$. Ein $f_{\nu}(r)$ hat den Ausdruck:

$$\frac{(r^{\nu} - 1) \left(r^{\frac{\nu}{\alpha\beta}} - 1 \right) \left(r^{\frac{\nu}{\alpha\gamma}} - 1 \right) \dots}{\left(r^{\frac{\nu}{\alpha}} - 1 \right) \left(r^{\frac{\nu}{\beta}} - 1 \right) \left(r^{\frac{\nu}{\gamma}} - 1 \right) \dots},$$

wenn $\alpha, \beta, \gamma, \dots$ die verschiedenen Primzahlen aus ν sind; also geht in $f_{\nu}(c)$ dieselbe Potenz von p auf wie in

$$\frac{p^{\nu} \cdot p^{\frac{\nu}{\alpha\beta}} \cdot p^{\frac{\nu}{\alpha\gamma}} \dots}{p^{\frac{\nu}{\alpha}} \cdot p^{\frac{\nu}{\beta}} \cdot p^{\frac{\nu}{\gamma}} \dots}$$

Diese Zahl ist 1, wenn ν sich aus mehreren ungleichen Primzahlen zusammensetzt, dagegen, wenn ν Potenz einer einzigen Primzahl ist, gleich

*) Hermite, Crelles Journal, Bd. 47, S. 312 (Oeuvres, T. I, p. 199); C. Jordan, Crelles Journal, Bd. 84, S. 112.

dieser Primzahl. Die höchste in $f_v(c)$ enthaltene Potenz von p ist demnach p^1 oder p^0 , je nachdem v eine Potenz von p ist oder nicht. Im ersteren Falle ist aber, $v > 1$ vorausgesetzt, $\varphi(v)$ mindestens gleich $p - 1 \geq 2$, im letzteren mindestens gleich 1.

Hat man nun $\Delta = (r - 1)^n$, so müssen, damit A eine endliche Ordnung besitze, auch alle $(n - 1)^{\text{ten}}$ Unterdeterminanten, also die Koeffizienten von Δ für $r = 1$ verschwinden, d. h. A ist die identische Transformation.

§ 2.

Es bezeichne f irgendeine positive quadratische Form mit n Variablen, reellen Koeffizienten und nichtverschwindender Determinante, und es sei $t(f)$ die Anzahl der verschiedenen ganzzahligen Transformationen, durch welche die Form in sich selbst übergeht.

Sollten in f die Koeffizienten nicht sämtlich in rationalen Verhältnissen zueinander stehen, so kann man immer leicht eine beliebig wenig von f verschiedene positive Form herstellen, in welcher solches der Fall ist, und welche dabei genau dieselben ganzzahligen Transformationen in sich zuläßt wie f . Letzteres tun ferner alle Formen, welche in den Verhältnissen ihrer Koeffizienten mit f ganz übereinstimmen. Im folgenden können wir deshalb voraussetzen, die Koeffizienten von f seien ganze Zahlen ohne gemeinsamen Teiler.

Die $t(f)$ Transformationen bilden eine Gruppe, und an anderer Stelle werde ich nachweisen, daß man, ausgehend von positiven quadratischen Formen, wenn auch nicht zu allen möglichen endlichen Gruppen ganzzahliger linearer Transformationen, so doch im besondern zu allen denjenigen gelangen kann, welche nicht in umfassenderen Gruppen als Untergruppen enthalten sind.

Die in Rede stehende Gruppe ist *einstufig isomorph* zur Gruppe der Reste ihrer Transformationen in bezug auf irgendeine ungerade Primzahl p . Denn lieferten zwei ihrer Transformationen, etwa A und B , gleiche Reste modulo p , so würde die Transformation $A^{-1} \cdot B$, welche, als Angehörige der Gruppe, ebenfalls von endlicher Ordnung wäre, der identischen Transformation modulo p kongruent, aber von ihr verschieden sein, was nach § 1 nicht angeht.

Die Gruppe der $t(f)$ Transformationenreste ist offenbar eine Untergruppe der Gruppe sämtlicher inkongruenter Transformationenreste modulo p von einer Determinante $\equiv \pm 1 \pmod{p}$, und ihre Ordnung, die Zahl $t(f)$, daher ein Divisor der Ordnung der letzteren Gruppe, d. i. der Zahl

$$(2) \quad 2(p^n - 1)p^{n-1}(p^{n-1} - 1)p^{n-2} \dots (p^2 - 1)p^*.$$

*) Galois, Journal de Liouville, T. XI, 1846, p. 410. (Oeuvres, p. 27.)

Jene Gruppe ist aber ebenso schon eine Untergruppe der Gruppe aller derjenigen Transformationenreste modulo p , welche die Form f modulo p ungeändert lassen. Die Ordnung dieser Gruppe hat für alle ungeraden Primzahlen p , welche nicht in der Determinante D der Form f aufgehen, also jedenfalls für *sämtliche* Primzahlen über einer gewissen Grenze l , den folgenden Ausdruck*), wenn n gerade ist:

$$(3) \quad p^{\frac{1}{4}n(n-2)} \cdot 2(p^2-1)(p^4-1)\dots(p^{n-2}-1)\left(p^{\frac{n}{2}}-\varepsilon\right),$$

wo $\varepsilon = \left(\frac{(-1)^{\frac{n}{2}}D}{p}\right)$ eine Einheit bedeutet; wenn n ungerade ist:

$$(4) \quad p^{\frac{1}{4}(n-1)^2} \cdot 2(p^2-1)(p^4-1)\dots(p^{n-1}-1).$$

Als Divisor *sämtlicher* Zahlen (3) oder (4) für ungerade Primzahlen $p > l$ ist die Zahl $t(f)$ auch ein Divisor des *größten gemeinschaftlichen Divisors* aller dieser Zahlen.

Um ein Resultat zu erhalten, das von der speziellen Form f unabhängig ist, denken wir uns in (3) den Faktor $p^{\frac{n}{2}} - \varepsilon$ durch sein Vielfaches $\frac{1}{2}(p^n - 1)$ ersetzt; ferner möge l mindestens gleich $n + 1$ sein. Dann ist jener größte gemeinsame Divisor dargestellt durch:

$$|\bar{n}| = \prod_q q^{\left[\frac{n}{q-1}\right] + \left[\frac{n}{q(q-1)}\right] + \left[\frac{n}{q^2(q-1)}\right] + \dots}, \quad (q = 2, 3, 5, 7, 11, \dots)$$

wenn unter der Bezeichnung $[a]$ die größte in a enthaltene ganze Zahl verstanden wird, und wenn q die Reihe der Primzahlen soweit durchläuft, bis das Produkt von selbst abbricht, d. i. bis zur größten Primzahl, welche noch $\leq n + 1$ ist.

Man hat, um dieses einzusehen, für eine jede Primzahl q eine Zahl (3) bzw. (4) aufzusuchen, welche eine möglichst niedrige Potenz von q enthält. Man gelangt zu einer solchen, indem man die Primzahl p in folgender Weise wählt: wenn $q > n + 1$ ist, als primitive Wurzel in bezug auf q ; wenn $q \leq n + 1$ und ungerade ist, als primitive Wurzel für den Modul q^2 ; wenn $q = 2$ ist, als Zahl der Form $\equiv \mp 1 + 4 \pmod{8}$. Die Existenz von Primzahlen p dieser Formen über der Grenze l ist eine Folge des bekannten Theorems über die arithmetischen Progressionen.

Im ersten der drei unterschiedenen Fälle ist dann die in Betracht kommende Zahl (3) oder (4) durch q überhaupt nicht teilbar; im zweiten

*) Vgl. *Untersuchungen über quadratische Formen*, Acta Mathematica, Bd. 7, S. 218. Diese Ges. Abhandlungen, Bd. I, S. 170.

geht q in $p^\nu - 1$ nur auf, wenn ν ein Vielfaches von $q - 1$ ist, und zwar dann in derselben Potenz wie in $q \cdot \frac{\nu}{q-1}$, mithin in der Zahl (3) bzw. (4) in derselben Potenz wie in $q^{\left[\frac{n}{q-1}\right]} \cdot 1 \cdot 2 \cdots \left[\frac{n}{q-1}\right]$; im dritten enthält $p^{2\nu} - 1$ dieselbe Potenz von 2 wie 8ν , also die Zahl (3) bzw. (4) dieselbe wie $2^{n + \left[\frac{n}{2}\right]} \cdot 1 \cdot 2 \cdots \left[\frac{n}{2}\right]$.

Die Identität der hier auftretenden Primzahlpotenzen mit denjenigen aus $\bar{n}|$ ergibt sich durch Anwendung der bekannten Relation:

$$1 \cdot 2 \cdots n = n! = \prod_q q^{\left[\frac{n}{q}\right] + \left[\frac{n}{q^2}\right] + \left[\frac{n}{q^3}\right] + \cdots}, \quad (q = 2, 3, 5, 7, \dots)$$

und man erhält damit in der Tat den Satz:

Die Anzahl der ganzzahligen Transformationen einer positiven quadratischen Form mit n -Variablen (und von nichtverschwindender Determinante) in sich selbst ist ein Divisor der Zahl $\bar{n}|$.

Als größter gemeinsamer Divisor aller Zahlen (2) würde sich $2^{\left[\frac{n}{2}\right]} \cdot \bar{n}|$ ergeben haben.

Die Zahl $\bar{n}|$ selbst ist ein Divisor von $(2n)!$; denn in ihrem Ausdrucke verkleinert man keinen der Exponenten, wenn man in denselben anstatt $q - 1$ überall $\frac{1}{2}q$ schreibt.

Als spezielle Fälle seien die folgenden erwähnt: die Form

$$\mathfrak{D}_n = x_1^2 + x_2^2 + \cdots + x_n^2$$

geht durch $2^n \cdot n!$, die Form

$$\mathfrak{D}_n = \left(\sum x_h \right)^2 + \sum x_h^2 *) \quad (h = 1, 2, \dots, n)$$

von der Determinante $n + 1$ durch $2 \cdot (n + 1)!$ ganzzahlige Transformationen in sich selbst über.

Die Zahl $\bar{n}|$ ist ferner das kleinste gemeinsame Vielfache aller möglichen Anzahlen $t(f)$.

Denn zunächst enthält die der Form $\mathfrak{D}_n$ angehörige Zahl $t(\mathfrak{D}_n)$ dieselbe Potenz von 2 wie $\bar{n}|$. Ist ferner q eine der ungeraden Primzahlen $\leq n + 1$, und bildet man eine Form f als Summe von $\left[\frac{n}{q-1}\right]$ Formen $\mathfrak{D}_{q-1}$ und der Form $\mathfrak{D}_{\left(n - (q-1)\left[\frac{n}{q-1}\right]\right)} = \mathfrak{D}$ mit fortlaufend nummerierten

*) Die Form $\mathfrak{D}_n$ ist von den Herren Korkine und Zolotareff eingehender untersucht worden, vgl. Mathematische Annalen, Bd. 6 und 11.

Variablen, so ist für diese Form f die Zahl

$$t(f) = (2 \cdot q!)^{\left[\frac{n}{q-1}\right]} \cdot \left[\frac{n}{q-1}\right]! t(\mathfrak{L})$$

durch dieselbe Potenz von q teilbar wie $\overline{n}$.

Man hat im einzelnen $\overline{2} = 24$, $\overline{3} = 48$, $\overline{4} = 5760$ usw. und allgemein:

$$\overline{2n+1} = 2 \cdot \overline{2n}, \quad \overline{2n} = 2b_n \cdot \overline{2n-1}, \quad \overline{n} = 2^n \cdot b_1 b_2 \dots b_{\left[\frac{n}{2}\right]},$$

wenn unter b_n eine Zahl verstanden wird, welche alle und nur solche Primzahlen q enthält, für welche $2n$ durch $q-1$ aufgeht, und jede derselben in einer Potenz q^{x+1} , falls sie in n in der Potenz q^x ($x \geq 0$) auftritt.

Bezeichnet B_n die n^{te} Bernoullische Zahl, so stellt b_n den Nenner von $\frac{1}{n} B_n$ vor*).

Die Bestimmung der ganzzahligen Transformationen einer positiven Form $f = \sum_1^n a_{hk} x_h x_k$ in sich selbst geschieht durch Vermittlung der äquivalenten reduzierten Formen.

Nach der Definition von Herrn Hermite**) gelten in einer positiven Klasse f diejenigen Formen als *reduziert*, welche das kleinste System

$$a_{11}, a_{22}, \dots, a_{nn}$$

(d. i. kurz ausgedrückt den kleinsten Wert von

$$a_{11} g^{n-1} + a_{22} g^{n-2} + \dots + a_{nn}$$

bei hinreichend großem positivem g) ergeben.

Den Sätzen, welche ich in Crelles Journal, Bd. 99 [[diese Ges. Abhandlungen, Bd. I, S. 153—156]] mitgeteilt habe, schließen sich die folgenden für den Fall von sechs Variablen an.

Eine Form f mit sechs Variablen ist immer und nur dann positiv und reduziert, wenn sie allen Ungleichungen

$$f(m_1, m_2, \dots, m_6) \geq a_{hh} \quad (h = 1, 2, \dots, 6)$$

genügt, für welche die Zahlen m in folgender Tabelle enthalten sind:

*) Vgl. Lipschitz, Crelles Journal, Bd. 96, S. 4.

**) Crelles Journal, Bd. 40, S. 302. (Oeuvres, T. I, p. 149.)

m_h	$\pm m_{h'}$	$\pm m_{h''}$	$\pm m_{h'''}$	$\pm m_{h^{IV}}$	$\pm m_{h^V}$
1	1				
1	1	1			
1	1	1	1		
1	1	1	1	1	
1	1	1	1	2	
1	1	1	1	1	1
1	1	1	1	1	2
1	1	1	1	2	2
1	1	1	1	2	3

(die nicht aufgeführten Größen m sind gleich Null zu setzen), und ferner den Ungleichungen:

$$a_{11} \leq a_{22} \leq \dots \leq a_{66}.$$

Nur für die reduzierten und die aus denselben durch Permutation der Variablen hervorgehenden Formen nehmen die Verbindungen

$$a_{11} + a_{22} + \dots + a_{66}, \dots, a_{11}a_{22} \dots a_{66}$$

ihre kleinsten Werte an.

Die Hermiteschen reduzierten Formen mit sieben Variablen lassen sich nicht mehr durch eine Reihe einzelner linearer Ungleichungen vollständig charakterisieren.

Berlin, den 15. Februar 1887.

VII.

Über die Bedingungen, unter welchen zwei quadratische Formen mit rationalen Koeffizienten ineinander rational transformiert werden können.

(Auszug aus einem von Herrn H. Minkowski in Bonn an Herrn Adolf Hurwitz gerichteten Briefe.)

(Crelles Journal für die reine und angewandte Mathematik, Band 106, S. 5—26.)

Bei unserem letzten Zusammensein interessierten Sie und Herr Hilbert sich für die Frage, unter welchen Umständen zwei diophantische Gleichungen zweiten Grades sich rational ineinander transformieren lassen*). In den folgenden Zeilen will ich eine Entscheidung dieser Frage zu geben versuchen.

Es liege irgendeine quadratische Form mit rationalen Koeffizienten vor, f . Die Anzahl ihrer Variablen heiße n , und bei dieser Variablenzahl habe die Form eine nichtverschwindende Determinante; als Aggregat von n Quadraten reeller linearer Formen dargestellt, möge f im ganzen I Quadrate mit negativem, und also $n - I$ mit positivem Vorzeichen aufweisen. Unterwirft man die Form einer beliebigen linearen Transformation mit rationalen Koeffizienten und von nichtverschwindender Determinante, so bleibt dabei zunächst außer den Zahlen n und I , da die Determinante

*) Bei der Untersuchung der ternären diophantischen Gleichungen vom Geschlechte Null wurden Herr Hilbert und ich auf diese Frage geführt. Wenn zwei diophantische Gleichungen durch rationale eindeutig umkehrbare Transformationen ineinander übergeführt werden können, so lassen sich offenbar die Lösungen einer jeden dieser Gleichungen aus den Lösungen der anderen ableiten; beide Gleichungen repräsentieren also im wesentlichen dieselbe Aufgabe. Wir rechnen deshalb alle diophantischen Gleichungen, welche aus einer durch die genannten Transformationen hervorgehen, in eine Klasse. Unsere Untersuchung, welche demnächst in den *Acta mathematica* erscheinen wird, [[*Acta mathematica* Bd. 14, S. 217—224]] ergab nun, daß in jeder Klasse ternärer diophantischer Gleichungen vom Geschlechte Null auch quadratische Gleichungen enthalten sind. Die sich hier anknüpfende Frage nach den Invarianten einer solchen Klasse findet daher durch die allgemeinen Sätze, welche Herr Minkowski aufstellt und beweist, ihre Erledigung.

A. Hurwitz.

der Form sich um das Quadrat der rationalen Transformationsdeterminante vervielfacht, die Gesamtheit aller der Primzahlen ungeändert, welche in der Determinante der Form in ungeraden Potenzen aufgehen. Das Produkt aller dieser Primzahlen, mit dem Vorzeichen $(-1)^I$ der Determinante genommen, heie A ; sind Primzahlen der bezeichneten Art nicht vorhanden, so werde $A = (-1)^I$ gesetzt.

Weiter werde ich nun zeigen, da, wenn p eine ganz beliebige Primzahl bedeutet, immer aus den Resten der Form f fr gengend hohe Potenzen von p als Moduln in gewisser Weise eine im allgemeinen nicht schon durch A allein bestimmte Gre sich herstellen lt — ich werde sie C_p nennen —, welche ihrer Bildung nach nur der Werte 1 oder -1 fhig ist, und welche den Wert, den sie fr f hat, bei keiner rationalen umkehrbaren Transformation von f verndert (vgl. unten Gleichung (1) und (2)). Fr alle diejenigen ungeraden Primzahlen, welche weder in der Determinante noch in dem Generalnenner der Koeffizienten von f wirklich vorkommen (d. h. in nichtverschwindenden Potenzen), findet sich diese Einheit von vornherein gleich $+1$, so da sie berhaupt nur fr eine endliche Anzahl von Primzahlen -1 sein kann. Diese Bemerkung vorweggenommen, besteht fr die Gesamtheit der Einheiten C_p von vornherein eine Beziehung zu den Werten n , I und A , nmlich ihr Produkt, also $C_2 C_3 C_5 C_7 C_{11} \dots$, erweist sich, wenn j die Anzahl der verschiedenen Primzahlen von der Form $4l + 3$ aus A bedeutet, gleich 1 oder -1 , je nachdem die Zahl $n - 2I - 2j$ modulo 8 den Rest 0, 1, 6, 7 oder 2, 3, 4, 5 lt. Auf diese Beziehung grnde ich eben den Nachweis fr die invariante Natur der Einheiten C_p : ich mache zuerst klar, da jedesmal bei einer Transformation hchstens diejenigen Einheiten C_p sich ndern knnten, welche den in der Determinante und dem Generalnenner der Transformation vorkommenden Primzahlen entsprechen, und dann nehme ich zu Hilfe, da eine jede rationale umkehrbare Transformation aus solchen besonderen zusammengesetzt werden kann, in deren Determinante und Generalnenner hchstens eine einzige Primzahl aufgeht. Da nun das Produkt aller Einheiten C_p invariant sein soll und deshalb sich nicht eine allein ndern kann, so bleiben bei derartigen Teiltransformationen und also auch bei jeder Transformation die Einheiten C_p ausnahmslos ungendert.

Im Falle $n = 2$ gelten auerdem noch folgende Beziehungen der Einheiten C_p zu dem von quadratischen Teilern befreiten Kern der Determinante: Sowie eine Primzahl p in A nicht vorkommt und $-A$ quadratischer Rest von p bzw., im Falle $p = 2$, von 8 ist, ist immer $C_p = 1$. — Im Falle $n = 1$ sind durch die Zahl A bereits smtliche Einheiten C_p bestimmt: Fr die in A nicht aufgehenden Primzahlen ist immer $C_p = 1$,

für die in A aufgehenden gleich 1 oder -1 , je nachdem $\frac{A}{p}$ quadratischer Rest oder Nichtrest von p ist bzw., im Falle $p=2$, die Form $8l \pm 1$ oder $8l \pm 3$ hat.

Nach dem, was über die Einheiten C_p bereits gesagt ist, müssen auch alle solchen ungeraden Primzahlen sich nach jeder rationalen Transformation von f beständig in der Determinante oder dem Generalnenner der transformierten Form wiederfinden, welche zwar der Zahl A nicht angehören, aber ein $C_p = -1$ ergeben; ist B das Produkt aller *dieser* Primzahlen, so wird daher, wenn die transformierte Form ganzzahlige Koeffizienten erhalten sollte, ihre Determinante den Faktor ABB haben müssen.

Durch eine Reihe von sehr einfachen ganzzahligen Transformationen mit der Determinante 1 und von solchen rationalen Transformationen, welche jede darin bestehen, daß sie eine einzelne Variable rational vervielfachen, gelingt es nun stets, wie ich zeigen werde, die vorgelegte Form in eine Form mit ganzzahligen Koeffizienten und genau von der Determinante ABB überzuführen. Dabei erweisen sich dann diejenigen arithmetischen Funktionen dieser Form, welche, wie man sich ausdrückt, ihr *Geschlecht* definieren, im allgemeinen als eindeutig durch I , A und die Einheiten C_p bestimmt. Nur, wenn zwischen n , A und dem Werte der Einheit C_2 eine gewisse Beziehung statthat, könnte noch die erhaltene Form zwei verschiedenen Geschlechtern von der Determinante ABB angehören; dann ist von diesen nur eines ungerade Zahlen darzustellen fähig, und sollte man nicht gerade zu einer Form dieses ungeraden Geschlechts gekommen sein, so kann man zu der erhaltenen Form stets solche äquivalente Formen finden, welche sich in Formen dieses Geschlechts in der einfachen Weise überführen lassen, daß man eine gewisse ihrer Variablen mit 2, eine gewisse andere mit $\frac{1}{2}$ multipliziert. Nun hat zuerst Henry John Stephen Smith*) ausgesprochen, daß irgend zwei Formen *eines* Geschlechts immer durch rationale Transformationen von der Determinante 1 und mit einem, zu einer beliebig vorgeschriebenen Zahl relativ primen Generalnenner ineinander übergeführt werden können, eine Eigenschaft, welche sie umgekehrt auch als zu demselben Geschlecht gehörig charakterisiert. Smith hat diesen fundamentalen Satz der Lehre von den quadratischen Formen in der Abhandlung „*On the Orders and Genera of Ternary Quadratic Forms*“ art. 12**) für ternäre Formen bewiesen, und auf Grund derselben Prinzipien und unter Zuhilfenahme gewisser Resultate aus meiner Arbeit „*Sur la théorie des formes quadratiques à coefficients*

*) Proceedings of the Royal Society of London, XVI. 1868. p. 202. (Collected Papers, vol. I, p. 516.)

**) Philosophical Transactions, CLVII. 1867. (Collected Papers, vol. I, p. 480.)

entiers“*) kann der Beweis für Formen mit beliebiger Variablenzahl geführt werden. Demnach wird es immer möglich sein, unsere Form f in eine bestimmte Form eines, durch die Einheiten C_p (und durch n und I) völlig bestimmten Geschlechts von der Determinante ABB rational zu transformieren.

Nun bezeichne B das Produkt aller überhaupt vorhandenen ungeraden Primzahlen, für welche $C_p = -1$ ist; dann sind aus dem Werte von B die Werte sämtlicher Einheiten C_p ersichtlich — der Wert von C_2 bei Benutzung der Relation für das Produkt aller $C_p =$, und man wird den Satz aussprechen dürfen:

Theorem I. Zwei rationale quadratische Formen mit n Variablen und nichtverschwindenden Determinanten können dann und nur dann rational ineinander transformiert werden, wenn sie gleiche Invarianten I , A und B haben.

Die vorher definierte Zahl B ist offenbar der Quotient aus B und dem größten Divisor von A und B .

Der Trägheitsindex I ist eine Zahl zwischen 0 bis n . A hat das Vorzeichen $(-1)^I$, B ist positiv; vom Vorzeichen abgesehen, sind A und B Produkte von lauter verschiedenen Primzahlen bzw. 1; B ist ungerade, in A kann unter Umständen die Primzahl 2 eingehen. — Im Falle $n=2$ enthält B niemals eine ungerade Primzahl, welche in A nicht vorkommt und von welcher — A quadratischer Rest ist, weil für jede solche Primzahl hier $C_p = 1$ ist, und wenn $-A \equiv 1 \pmod{8}$ ist, ist immer $C_2 = 1$ und muß infolgedessen die Anzahl der in B vorkommenden Primzahlen mit der alsdann offenbar ganzzahligen Größe $\frac{1}{2}(1 - I - j)$ zugleich gerade oder ungerade ausfallen; dabei soll j wieder die Anzahl der Primzahlen von der Form $4l + 3$ aus A vorstellen. — Im Falle $n = 1$ ist B das Produkt aller derjenigen ungeraden Primzahlen p aus A , für welche $\frac{A}{p}$ quadratischer Nichtrest von p ist.

Zu jedem mit den vorstehenden Bedingungen verträglichen Systeme von Zahlen n , I , A , B gibt es wirklich ein oder zwei Geschlechter ganzzahliger Formen von der zugehörigen Determinante ABB , und können Repräsentanten dieser Geschlechter nach pp. 89—90 meiner vorher zitierten Arbeit [[diese Ges. Abhandlungen, Bd. I, S. 76—77]] aufgestellt werden.

Sie sehen hiernach, daß bei gegebener Variablenzahl und gegebenem Trägheitsindex stets unendlich viele verschiedene Typen von quadratischen Formen, die nicht rational ineinander zu transformieren sind, existieren, und daß von $n = 3$ an die unendliche Anzahl dieser Typen in gewissem

*) Mémoires présentés à l'Académie des Sciences de l'Institut de France, XXIX. Nr. 2. 1884. Diese Ges. Abhandlungen, Bd. I, S. 3—144.

Sinne durch eine Potenz von 2 ausgedrückt werden kann, deren Exponent die um 1 verminderte doppelte Anzahl aller natürlichen Primzahlen ist.

Man kann die Invarianten n, I, A, B einer zerfallenden quadratischen Form, d. h. einer solchen, welche als Summe zweier Formen mit verschiedenen Variablen erscheint, stets durch die entsprechenden Invarianten ihrer Teile und umgekehrt die eines Teiles durch die der Form und des anderen Teiles darstellen. Nämlich die Zahlen n der Teile und ebenso die Zahlen I der Teile geben summiert die entsprechende Zahl der Summe; die Invariante A der Summe ist, von ihrem Vorzeichen $(-1)^I$ abgesehen, das Produkt aller der Primzahlen, welche nur in einer der Zahlen A der zwei Teile vorkommen, und endlich ist (vgl. unten Gleichung (3)) die Zahl B der Summe das Produkt aller der Primzahlen von der Form $4l + 1$, welche nur in einer der Zahlen B der Teile vorkommen, und aller der Primzahlen von der Form $4l + 3$, welche entweder in beiden A und in beiden oder in keinem der B der Teile, oder in einem der B und zugleich in einem oder in keinem der A der Teile vorkommen.

Hält man dieses Resultat mit dem Umstande zusammen, daß von $n = 3$ an die Invarianten A und B völlig unabhängig voneinander sind, so ist insbesondere zu schließen, daß man zu jeder Form f mit mehr als drei Variablen nach Belieben eine Form g mit drei Variablen weniger annehmen kann, welche nur, als Aggregat von soviel Quadraten reeller linearer Formen dargestellt, als sie Variablen besitzt, weder mehr positive noch mehr negative Quadrate enthalten darf als f in einer entsprechenden Darstellung, und daß dann jedesmal ein Typus von rationalen Formen χ mit drei Variablen existiert, so daß f rational in $g + \chi$ transformiert werden kann. Beispielsweise kann $g = \sum \pm x_h^2 (h = 1, 2, \dots, n - 3)$ genommen werden, wenn darin nicht mehr als I Koeffizienten -1 und nicht mehr als $n - I$ gleich 1 sind, wobei n und I sich auf f beziehen. —

Bei der von Ihnen angeregten Frage nun handelt es sich nicht darum, ob zwei quadratische Formen sich direkt rational ineinander transformieren lassen, sondern ob vielleicht eine solche Transformation dadurch möglich gemacht werden kann, daß man die Formen noch mit geeigneten Faktoren versieht. Es ist also noch zu untersuchen, in welchen Punkten die Zahlen n, I, A, B der verschiedenen rationalen Multipla einer und derselben Form f alle miteinander übereinstimmen und in welchen nicht.

Je nachdem der Faktor, mit welchem man eine Form f multipliziert, positiv oder negativ ist, ändert sich ihr Index I nicht, oder er geht in $n - I$ über; in jedem Falle bleibt der absolute Wert von $n - 2I$ unverändert; ich bemerke, daß $n - 2I$ die Differenz zwischen der Anzahl der Quadrate mit positivem und der mit negativem Vorzeichen in einer Darstellung von f als Aggregat von n Quadraten reeller linearer Formen ist. Wenn n

gerade und $I = \frac{n}{2}$ ist, so behält der Index selbst in jedem Falle seinen Wert.

Wir können uns hier auf Hinzufügung solcher rationaler Faktoren M zu f beschränken, die, vom Vorzeichen abgesehen, Produkte von lauter verschiedenen Primzahlen sind, indem wir uns jeden rationalen quadratischen Faktor mit den Variablen der Form verschmolzen denken können. Die Determinante von f nimmt in Mf den Faktor M^n an. Wenn n gerade ist, so bleibt also ihr von quadratischen Teilern entblößter Kern A ungeändert; ist n ungerade, so gibt es unter allen diesen Vielfachen Mf ein einziges, dessen Determinantenkern 1 ist, nämlich die Form Af , und es wird dann die dieser Form zukommende Zahl B als Invariante der diophantischen Gleichung $f = 0$ zu betrachten sein.

Um die hinsichtlich der Einheiten C_p obwaltenden Verhältnisse klarzulegen, möge mit $\delta(p)$ oder δ die Zahl 0 bzw. 1 bezeichnet werden, je nachdem die gerade betrachtete Primzahl p in der Zahl A von f nicht vorkommt bzw. darin vorkommt. Unter c_p soll dann die positive oder negative Einheit verstanden werden, je nachdem die, p nicht mehr enthaltende Zahl $(-1)^{\left[\frac{n}{2}\right] + (n-1)\delta} \cdot p^{-\delta} A$ quadratischer Rest oder Nichtrest von p ist bzw., falls $p = 2$ ist, die Form $8l \pm 1$ oder $8l \pm 3$ hat; endlich soll c ohne Index diejenige Einheit vorstellen, welcher die ungerade Zahl $(-1)^{\left[\frac{n}{2}\right]} 2^{-\delta(2)} A$ modulo 4 kongruent ist; $\left[\frac{n}{2}\right]$ bedeutet hier immer die größte in $\frac{n}{2}$ enthaltene ganze Zahl. Wenn n gerade ist, so sind für alle Formen Mf die Zahlen δ dieselben. Aus den später für die Einheiten C_p aufzustellenden Gleichungen (vgl. unten (1) und (2)) wird nun unmittelbar zu schließen sein:

Von der Einheit C_p einer Form f unterscheidet sich die entsprechende Einheit einer Form Mf , wenn $M = p$ ist, um den Faktor c_p ; wenn M zu p relativ prim ist, um den Faktor $\left(\frac{M}{p^\delta}\right)$ oder $\left(\frac{c^{n-1}2^\delta}{M}\right)$, je nachdem p ungerade oder 2 ist. Die Klammern bedeuten das Legendre-Jacobische Symbol, und zwar setze ich dasselbe hier für den Fall negativer Zahlen im Zähler wie im Nenner durch die Gleichung $\left(\frac{Z}{N}\right) = -\left(\frac{Z}{-N}\right)$ definiert voraus, so daß insbesondere für alle ungeraden Zahlen N , für positive sowohl wie für negative, das Symbol $\left(\frac{-1}{N}\right) = (-1)^{\frac{N-1}{2}}$ ist. Hieraus folgt nun:

1) Wenn n gerade ist, so hat eine Einheit C_p dann und nur dann für alle Formen Mf gleichen Wert, wenn p in A nicht vorkommt ($\delta = 0$) und $(-1)^{\frac{n}{2}} A$ quadratischer Rest von p bzw., wenn $p = 2$, von 8 ist ($c_p = 1$

bzw. $c = 1$ und $c_2 = 1$). Bezeichnet also B_1 das Produkt aller derjenigen von den hier in Betracht kommenden Primzahlen, für welche $C_p = -1$ ist, so stellt dieses B_1 , ebenso wie A , hier eine Invariante der Gleichung $f = 0$ vor. B_1 ist, von einem eventuellen Faktor 2 abgesehen, ein Divisor der früher definierten Zahl B ; multipliziert man f mit allen außerdem noch in B vorkommenden ungeraden Primzahlen (d. h. für die $\delta = 0$, $c_p = -1$, $C_p = -1$ ist) und noch mit der Primzahl 2, falls $(-1)^{\frac{n}{2}} A \equiv 5 \pmod{8}$ (d. i. $\delta(2) = 0$, $c = 1$, $c_2 = -1$) und $C_2 = -1$ ist, so entsteht eine Form, die $B_2 f$ heißen möge, welche unter den Vielfachen Mf die Eigenschaft besitzt, für möglichst wenige von den in A nicht vorkommenden ungeraden Primzahlen, nämlich nur für diejenigen aus B_1 , ein $C_p = -1$ zu ergeben, und auch, wenn $(-1)^{\frac{n}{2}} A \equiv 1 \pmod{4}$ ist, insofern dies überhaupt angeht, nicht ein $C_2 = -1$ zu liefern.

Da die Zahlen A und B_1 Produkte von lauter verschiedenen Primzahlen sind, so sind die Werte dieser beiden Zahlen aus dem Werte der einen Zahl $D = AB_1 B_1$ zu entnehmen, welche, da B_1 zu A relativ prim ist, keine Primzahl in einer höheren als der zweiten Potenz enthalten wird. Der Wert dieser Zahl D ist außer durch das Verhältnis, in welchem die einzelnen Primzahlen von B_1 zu A stehen sollen, durch die bereits oben erwähnten, für die Einheiten C_p von vornherein gegebenen Beziehungen in der Weise beschränkt, daß, wenn $(-1)^{\frac{n}{2}} A = 1$ (d. h. $A = \pm 1$ und $\frac{1}{2}(n - 2I)$ gerade) ist, die Anzahl der in B_1 auftretenden Primzahlen zugleich mit der Zahl $\frac{1}{4}(n - 2I)$ gerade oder ungerade sein muß, und daß im Falle $n = 2$ immer $B_1 = 1$ ist.

2) Wenn n ungerade ist, so möge der Buchstabe D zur Bezeichnung der Invariante B der Form Af verwandt werden; diese findet sich dann, durch die Invarianten von f ausgedrückt, gleich $A_1 B$, wenn A_1 das Produkt aller derjenigen ungeraden Primzahlen aus A bedeutet, für welche $C_p = -c_p$ ist. Im Falle $n = 1$ ist immer $D = 1$.

Theorem II. Wenn zwei Formen mit n Variablen und mit nichtverschwindenden Determinanten in den absoluten Werten ihrer Zahlen $n - 2I$ und in ihren Invarianten D übereinstimmen, so kann jede von ihnen rational in ein rationales Vielfaches der anderen transformiert werden.

Ich beweise diesen Satz, indem ich ihn auf den Satz I zurückführe:

1) Ist n gerade, so multipliziere man jede der beiden Formen mit ihrer Zahl B_2 , und falls die Indizes der Formen nicht gleich sind (also sich zu n ergänzen und beziehungsweise unter und über $\frac{1}{2}n$ liegen), noch eine der Formen mit -1 . Man hat dann zwei Formen, φ , ψ , welche gleichen Index und dasselbe A besitzen, und für alle in ihrer Zahl A nicht

aufgehenden ungeraden Primzahlen gleiche Einheiten C_p liefern, und überdies, wenn $(-1)^{\frac{n}{2}} A \equiv 1 \pmod{4}$ (d. h. $\delta(2) = 0$, $c = 1$) ist, auch dasselbe C_2 ergeben; und es kommt darauf an, eine der Formen, etwa ψ , noch mit einem solchen positiven Faktor zu multiplizieren, daß eine Form entsteht, welche mit der anderen Form φ in *allen* Einheiten C_p übereinstimmt. Ist nun M irgendeine zu $2D$ relativ prime, positive ganze Zahl, welche den Bedingungen genügt, daß für jede in A aufgehende ungerade Primzahl p $\left(\frac{M}{p}\right) = 1$ oder $= -1$ ist, je nachdem die Charaktere C_p für φ und ψ gleich oder verschieden sind, und daß, wenn *nicht* $(-1)^{\frac{n}{2}} A \equiv 1 \pmod{4}$ ist, das Symbol $\left(\frac{c2^{\delta(2)}}{M}\right) = 1$ oder $= -1$ ist, je nachdem φ und ψ gleiche oder verschiedene Charaktere C_2 haben, so kann bei $M\psi$ und φ ein Zweifel über die Gleichheit der Einheiten C_p nur noch hinsichtlich der in M aufgehenden Primzahlen möglich sein. Die genannten Bedingungen für M sind aber derart, daß man ihnen durch eine *Primzahl* gerecht werden kann, in welchem Falle es sich nur noch um den einen Charakter C_M handeln wird; und dieser wird dann für $M\psi$ und φ deshalb nicht verschieden sein können, weil er sich für beide Formen in gleicher Weise durch das Produkt der übrigen Charaktere C_p ausdrücken läßt.

2) Ist n ungerade, so multipliziere man jede der vorgelegten Formen mit ihrer Zahl A , und man erhält so zwei Formen, für die alle Bedingungen dafür erfüllt sind, daß sie sich rational ineinander transformieren lassen.

Es mögen noch einige spezielle Folgerungen aus den Sätzen I und II Erwähnung finden.

Eine rationale Form kann dann und nur dann in irgendwelche *negative*, rationale Multipla von sich selbst rational transformiert werden, wenn ihre Variablenzahl n gerade und ihr Trägheitsindex gleich $\frac{1}{2}n$ ist.

Eine Form f ist dann und nur dann in $-f$ rational zu transformieren, wenn ihre Variablenzahl n gerade, ihr Trägheitsindex gleich $\frac{1}{2}n$ ist und ihre Determinante keine Primzahl von der Form $4l + 3$ in ungerader Potenz enthält.

Durch welche quadratische Formen mit n Variablen können von Null verschiedene, rationale Quadratzahlen dargestellt werden? Ist durch irgendeine Form eine von Null verschiedene Zahl N darstellbar, so läßt die Form sich stets so rational transformieren, daß sie die Zahl N als ersten Koeffizienten erhält, und dann geht sie bei dem ersten Schritte der gewöhnlichen Methode, eine quadratische Form in Einzelformen mit einer Variable zu zerlegen, über in $Nx^2 +$ einer Form, welche die Variable x nicht mehr enthält. Es handelt sich also hier um die Kriterien für solche Formen mit n Variablen, welche sich rational transformieren lassen in

das Quadrat einer Variable + einer Form mit $n - 1$ Variablen, und man findet:

Von Null verschiedene rationale Quadratzahlen sind darstellbar durch jede nicht wesentlich negative Form mit 4 oder mehr Variablen; durch jede nicht wesentlich negative Form mit drei Variablen, deren Invarianten A und B auch bei Formen mit zwei Variablen vorkommen können (die hierfür zu erfüllenden Bedingungen s. oben); durch jede nicht wesentlich negative Form mit zwei Variablen, deren Invarianten A und B auch bei einer Form mit einer Variable vorkommen können.

Wann kann durch eine rationale Form f die Zahl Null rational dargestellt werden (natürlich, ohne daß alle Variablen verschwindende Werte erhalten)? Es sei f irgendwie in eine Summe von n Formen mit einer Variablen transformiert, und in einer Lösung von $f = 0$ etwa die Variable einer dieser Formen, welche Nx^2 heißen möge, von Null verschieden und $= x_1$. Die Form f besteht dann aus Nx^2 und einer Form mit $n - 1$ Variablen; durch letztere wird, während durch f die Zahl Null dargestellt wird, die Zahl $-Nx_1^2$ dargestellt, und diese Form läßt sich daher nach dem vorher Bemerkten rational transformieren in eine Form $-Ny^2$ + einer Form mit $n - 2$ Variablen. Da nun $N(x^2 - y^2)$ stets in $x^2 - y^2$ rational zu transformieren ist, so muß also f in $x^2 - y^2$ + einer Form mit $n - 2$ Variablen zu transformieren sein, wenn $f = 0$ lösbar sein soll; so ergibt sich:

Die Zahl Null ist rational darstellbar

*durch jede indefinite Form mit fünf oder mehr Variablen,
durch jede indefinite Form mit vier Variablen, deren Invariante D
nur erste (keine zweiten) Potenzen von Primzahlen enthält,
durch jede indefinite Form mit drei Variablen, deren Invariante D
gleich 1 ist,
durch jede (indefinite) Form mit zwei Variablen, deren Invariante
D = -1 ist.*

Betreffs der über die Gleichung $ax^2 + by^2 + cz^2 = 0$ existierenden Literatur sei auf die Vorlesungen über Zahlentheorie von Dirichlet, herausgegeben von Dedekind, III. Aufl. Suppl. X [[IV. Aufl. Suppl. X, S. 418 ff.]] verwiesen. Kriterien für die Darstellbarkeit der Zahl Null durch beliebige ternäre Formen sind von H. J. S. Smith*) ohne Beweis publiziert: dieselben sind dann später von Herrn A. Meyer**) begründet worden. Ferner hat Herr Meyer***) die notwendigen und hinreichenden Bedingungen für

*) Proceedings of the Royal Society of London, XIII, p. 110. (Collected Papers, vol. I, p. 410.)

**) Crelles Journal für Mathematik, Bd. 98, S. 177.

***) Mathematische Mitteilungen, Vierteljahrsschrift der naturforschenden Gesellschaft in Zürich, XXIX. S. 209.

die Auflösbarkeit der Gleichung $ax^2 + by^2 + cz^2 + du^2 = 0$, doch in einer umständlicheren Fassung als der hier gegebenen, aufgestellt und bewiesen, daß durch eine jede indefinite Form mit mehr als vier Variablen die Zahl Null rational darstellbar ist. — Endlich ist zu erwähnen, daß Eisenstein*) die Frage behandelt hat, welche ternären Formen sich rational in ein Multiplum von $x^2 + y^2 + z^2$ transformieren lassen.

Ich werde nun zunächst das Bildungsgesetz der Einheiten C_p ableiten.

Jeder rationalen quadratischen Form $f(x_1, x_2, \dots, x_n)$ mit gebrochenen Koeffizienten kann in ganz bestimmter Weise eine aus ihr durch eine rationale Transformation abzuleitende Form mit ganzen Koeffizienten zugeordnet werden, nämlich, wenn N den Generalnenner der Koeffizienten von f bedeutet, die Form NNf , in welche f durch die Substitution $x_h = Ny_h$ ($h = 1, 2, \dots, n$) übergeht. Wir können uns deshalb auf die Betrachtung ganzzahliger Formen beschränken, und brauchen mit denselben auch nur solche Transformationen vorzunehmen, durch welche wir wieder auf ganzzahlige Formen kommen.

Es ist klar, daß wir solche Funktionen einer quadratischen Form, welche für alle rationalen umkehrbaren Transformationen der Form invariant sein sollen, nur unter den, das Geschlecht der Form definierenden Größen zu suchen haben; Größen nämlich, die nicht einmal für alle Formen eines Geschlechts gleichwertig sind, erleiden immer schon bei gewissen rationalen Transformationen von der Determinante 1 Änderungen.

Die Invarianten des Geschlechts einer quadratischen Form sind ihre *Ordnung* und ihre *Charaktere*. (Die nun zunächst folgenden Bemerkungen sind meinem Aufsätze „Sur la théorie des formes quadratiques“ entnommen, den ich weiterhin kurz mit *F. Q.* zitieren will.**) Heißt die Form $f = \sum_1^n a_{hk} x_h x_k$ ($a_{hk} = a_{kh}$), und sind alle Koeffizienten ganze Zahlen und die Determinante $|a_{hk}| = \Delta$ von Null verschieden, so rechnet man zu den Bestimmungsstücken der Ordnung der Form ihre Zahl n , ihren Trägheitsindex I , ihre Determinante Δ , ferner die größten gemeinsamen Teiler aller einreihigen, aller zweireihigen usw. bis aller $(n - 1)$ -reihigen Unterdeterminanten von $|a_{hk}|$ — diese Teiler mögen der Reihe nach $d_0, d_1, \dots, d_{n-2}$ heißen, endlich diejenigen größten gemeinsamen Teiler, welche an die Stelle dieser Teiler treten, wenn von den betreffenden Unterdeterminanten die unsymmetrischen mit dem Faktor 2 multipliziert genommen werden

*) Journal de Liouville. XVII. 1852. S. 473.

**) In [[]] fügen wir diesen Zitaten den entsprechenden Hinweis auf die in diesen Gesammelten Abhandlungen, Bd. I, S. 3—144, veröffentlichte deutsche Ausgabe hinzu. (Anm. d. Herausg.)

— diese letzteren Teiler bezeichne ich mit $\sigma_1 d_0, \sigma_2 d_1, \dots, \sigma_{n-1} d_{n-2}$, so daß eine jede Zahl σ_h gleich 1 oder gleich 2 ist; ich setze noch $\Delta = (-1)^t d_{n-1}$ und $\sigma_0 = 1, \sigma_n = 1$. Die durch die Gleichungen

$$d_h = d_0^{h+1} o_1^h o_2^{h-1} \dots o_h \quad (h=1, 2, \dots, n-1)$$

bestimmten $n-1$ Größen o_h erweisen sich immer als ganze Zahlen (*F. Q.*, p. 6 [[S. 13]]), und ferner sind die Zahlen σ_h immer so beschaffen, daß ein Produkt $\sigma_{h-1} o_h \sigma_{h+1}$ ungerade ist, sowie $\sigma_h = 2$ ist, und durch 4 teilbar, sowie $\sigma_{h-1} = 2$ oder $\sigma_{h+1} = 2$ ist (*F. Q.*, p. 31 [[S. 31]]).

Die Charaktere sind Größen, die in Gestalt Legendrescher Symbole auftreten, also Einheiten ± 1 ; ihre Werte können, wie ich *F. Q.* artt. VII-IX [[S. 45—70]] gezeigt habe, ausnahmslos erschlossen werden aus den Werten der verschiedenen Summen*)

$$f(\alpha, N) = \sum e^{\frac{2\pi i \alpha f(x_1, x_2, \dots, x_n)}{N}} \quad \begin{pmatrix} x_h = 1, 2, \dots, N \\ h = 1, 2, \dots, n \end{pmatrix},$$

in welchen N und α zueinander relativ prime ganze Zahlen bedeuten und die Variablen Restsysteme modulo N zu durchlaufen haben; e ist hier die Basis der natürlichen Logarithmen, π die Ludolphsche Zahl, $i = \sqrt{-1}$. Die Werte solcher Summen $f(\alpha, N)$ hängen offenbar nur von den Resten von f in bezug auf die Moduln N ab.

Indem man die Betrachtungen, welche *F. Q.*, p. 60 beziehungsweise p. 65 [[S. 54 und 58]] abgebrochen sind, fortführt, gelangt man in betreff dieser Summen insbesondere zu folgendem wichtigen Resultate:

Zu jeder Primzahl p gehört eine gewisse Einheit C_p von solcher Art, daß für jede Potenz p^t , welche nicht niedriger als die in $\frac{4\Delta}{\sigma_{n-1} d_{n-2}}$ aufgehende Potenz von p ist, und jede beliebige zu p relativ prime Zahl α , wenn noch p^∂ die höchste in Δ enthaltene Potenz von p bedeutet und die Einheiten c_p und c in der bereits oben festgesetzten Beziehung zur Determinante stehen, die Gleichung gilt: wenn p ungerade ist,

$$(1) \quad f(\alpha, p^t) = \left(\frac{\alpha}{p^{\partial+n t}} \right) C_p c_p^t i^{\left(\frac{p^{\partial+n t}-1}{2} \right)^2} p^{\frac{\partial+n t}{2}};$$

wenn $p = 2$ ist,

$$(2) \quad f(\alpha, 2^t) = \left(\frac{c^{n-1} 2^{\partial+n t}}{\alpha} \right) C_2 c_2^t (-i)^{\left(\frac{\alpha^n c - 1}{2} \right)^2} \left(\frac{1+i}{\sqrt{2}} \right)^{n-2} \left[\frac{n}{2} \right] 2^{\frac{\partial+n t+n}{2}}.$$

Die erste Klammer auf der rechten Seite ist jedesmal ein Legendresches Symbol. Ich bemerke noch, daß $\sigma_{n-1} d_{n-2}$ den größten gemeinsamen Teiler der sämtlichen Zahlen $\frac{\partial \Delta}{\partial \alpha_{h h}}$ und $2 \frac{\partial \Delta}{\partial \alpha_{h k}}$ ($h \neq k$) vorstellt.

*) Diese Summen bilden auch den Gegenstand des Aufsatzes von Herrn H. Weber: Ueber die mehrfachen Gaußschen Summen, *Crelles Journal*, Bd. 74, S. 214—256.

Da für solche ungerade Primzahlen, welche in der Determinante Δ nicht vorkommen, der Exponent t in (1) bereits gleich Null genommen werden darf, so ist für alle diese Primzahlen offenbar $C_p = 1$. Überhaupt haben wir hier diejenigen Einheiten vor uns, von welchen oben die Rede gewesen ist, wenn wir nur noch festsetzen, daß unter den Einheiten C_p einer Form f mit einem Generalnenner $N > 1$ die entsprechenden Einheiten der ganzzahligen Form NNf verstanden werden sollen.

Daß eine jede der vorstehend definierten Einheiten C_p ungeändert bleibt bei allen solchen rationalen umkehrbaren Transformationen von f , bei welchen Determinante und Generalnenner zu der betreffenden Primzahl p relativ prim sind, und welche aus f wieder ganzzahlige Formen hervorgehen lassen, ist ohne weiteres ersichtlich. Denn da durch solche Transformationen aus zwei, für einen Modul p^t inkongruenten Systemen der Variablen von f immer wieder inkongruente Systeme der neuen Variablen, und also aus vollständigen Restsystemen in bezug auf einen Modul p^t wieder solche Systeme entstehen, so erfahren dabei die Summen $f(\alpha, p^t)$ keine Veränderung, und da offenbar auch die Potenz p^d sich nicht ändert, so behält auch die Einheit C_p ihren Wert bei.

Sehen wir nun zu, wie diese Einheiten C_p zu berechnen sind. Zerfällt eine Form f , in bezug auf einen Modul p^t betrachtet, in die Summe zweier Formen, f' , f'' , mit verschiedenen Variablen, so ist jede ihr zugehörige Summe $f(\alpha, p^t)$ das Produkt der analogen Summen für f' und f'' , und so folgt aus (1) für ein ungerades p :

$$(3) \quad C_p = (-1)^{\frac{p^{\partial'}-1}{2} \frac{p^{\partial''}-1}{2}} C_p' C_p'',$$

und aus (2) für $p = 2$:

$$(4) \quad C_2 = (-1)^{n' n'' \frac{c-1}{2} + \frac{c'-1}{2} \frac{c''-1}{2}} C_2' C_2'';$$

die gestrichelten Buchstaben sind auf die Formen f' und f'' zu beziehen.

Nun kann man jede Form f , wenn von ihren Zahlen $o_1, o_2, \dots, o_{n-1}$ im ganzen $\lambda - 1$ durch eine Primzahl p teilbar sind, durch höchstens $n - \lambda$ Substitutionen, welche darin bestehen, daß sie eine Variable um eine zweite vermehren, und höchstens $n - 1$ Substitutionen, welche darin bestehen, daß sie zwei Variablen miteinander permutieren und gleichzeitig eine derselben mit -1 multiplizieren, in eine solche Form φ transformieren, in welcher die aus den ersten $h = 1, 2, \dots, n - 1$ Horizontal- und Vertikalreihen der Determinante gebildeten Unterdeterminanten möglichst niedrige Potenzen von p als Faktoren haben, d. h. Werte $\sigma_h d_{h-1} \varphi_h$ besitzen, in denen die Zahlen φ_h zu p relativ prim sind (*F. Q.*, p. 36 [[S. 35]]); es werde noch $\varphi_0 = 1$, $\varphi_n = (-1)^t$ gesetzt. Eine derartige Form φ nenne ich eine *Grundform* in bezug auf p , und eine Grundform kann weiter für

jeden Modul p^t mit Hilfe von solchen Substitutionen, in deren Determinante alle Glieder links von der Diagonale Null und alle Glieder in der Diagonale 1 sind und welche daher die Werte der Zahlen $\sigma_h d_{h-1} \varphi_h$ ($h = 1, 2, \dots, n$) sämtlich ungeändert lassen, in sogenannte *Hauptreste* umgewandelt, d. h. dergestalt transformiert werden, daß sie modulo p^t in eine Summe von Formen mit einer Variable oder, wenn $p = 2$ ist, in Formen mit einer Variable und Formen zweiter Art mit zwei Variablen zerfällt. Bringt man dann die Formeln (3) und (4) in Anwendung und benutzt zugleich die bereits oben auseinandergesetzten und nun aus (1) und (2) fast unmittelbar zu entnehmenden Relationen, welche die Einheiten C_p zweier Formen verbinden, die rationale Vielfache voneinander sind, so wird man in letzter Instanz auf die Einheiten C_p der Form xx geführt, die sich sämtlich gleich 1 erweisen (vgl. Gauß, *Summatio quarumdam serierum singularium*, Werke, Bd. II), und auf Einheiten C_2 für Formen $\alpha xx + \beta xy + \gamma yy$ mit ganzen α, β, γ und ungeraden β , die ebenfalls gleich 1 sind (F. Q., pp. 57–59 [[S. 51–53]]).

Die Exponenten der in den Zahlen $d_0, o_1, o_2, \dots, o_{n-1}$ der Form f aufgehenden Potenzen von p mögen $v_0, \omega_1, \omega_2, \dots, \omega_{n-1}$ heißen, und es werde noch gesetzt $v_0 + \omega_1 + \omega_2 + \dots + \omega_h = v_h$ ($h = 1, 2, \dots, n-1$); ferner mögen mit D_h ($h = 1, 2, \dots, n$) diejenigen, zu p relativ primen Zahlen bezeichnet werden, die sich ergeben, wenn die Zahlen $\sigma_h d_{h-1} \varphi_h$ einer mit f äquivalenten Grundform φ in bezug auf p von ihren Potenzen von p befreit werden. Die aus einer solchen Grundform entstehenden Hauptreste lauten dann, wenn p ungerade ist (F. Q., p. 34 [[S. 34]]):

$$(5) \quad p^{v_0} D_1 x_1^2 + p^{v_1} \frac{D_2}{D_1} x_2^2 + \dots + p^{v_{n-1}} \frac{D_n}{D_{n-1}} x_n^2 \pmod{p^t},$$

und man findet auf dem soeben angegebenen Wege:

$$C_p = \left(\frac{-1}{p^{\sum v_h \varphi_k}} \right) \left(\frac{D_n}{p^{v_{n-1}}} \right) \prod \left(\frac{D_h}{p^{\omega_h}} \right). \quad \begin{matrix} (h=1, 2, \dots, n-1) \\ (k=0, 1, \dots, h-1) \end{matrix}$$

Wenn $p = 2$ ist, so können die aus einer Grundform entspringenden Hauptreste so geschrieben werden:

$$(6) \quad \sum_{(\sigma_h=1, \sigma_{h-1}=1)} \left\{ 2^{v_h-1} \frac{D_h}{D_{h-1}} x_h^2 \right. \\ \left. \text{bzw. } 2^{v_h-1+1} \left(\frac{D_h}{D_{h-1}} x_h^2 + u_h x_h x_{h+1} + \frac{D_{h+1} + u_h^2 D_{h-1}}{4 D_h} x_{h+1}^2 \right) \right\} \pmod{2^t}, \\ (\sigma_h=2; \sigma_{h-1}=1, \sigma_{h+1}=1)$$

worin die binären Formen *denjenigen* Größen σ_h entsprechen, welche gleich 2 sind (und welche immer die Relationen $\sigma_{h-1} = 1, \sigma_{h+1} = 1$ und $-D_{h-1} \equiv D_{h+1} \pmod{4}$ im Gefolge haben), und wobei die u_h in diesen binären Formen beliebige ungerade Zahlen bedeuten; man erhält dann:

$$C_2 = (-1)^{\left[\frac{n}{4}\right] + \left[\frac{n}{2}\right] \left(\left[\frac{n}{2}\right] + \frac{D_n - 1}{2}\right)} (-1)^{\sum \frac{D_h - 1}{2} \frac{D_{h+1} + 1}{2}} \left(\frac{\sigma_{n-1} 2^{v_{n-1}}}{D_n}\right) \prod \left(\frac{\sigma_{h-1} 2^{v_h} \sigma_{h+1}}{D_h}\right). \\ (h = 1, 2, \dots, n-1)$$

Für die Gesamtheit dieser Einheiten C_p besteht eine Beziehung zur Determinante von f , die sich folgendermaßen herleiten läßt. Man kann zu f immer solche äquivalente Formen φ finden, welche sowohl für die Primzahl 2 wie für alle in der Determinante von f aufgehenden ungeraden Primzahlen Grundformen sind und in welchen außerdem je zwei aufeinanderfolgende der Zahlen $\varphi_1, \varphi_2, \dots, \varphi_{n-1}$ zueinander relativ prim sind (F. Q., p. 83 [[S. 72]]). Für solche Formen φ führen die quadratischen Kongruenzen:

$$-\sigma_{h-1} \sigma_h \sigma_{h+1} \varphi_{h-1} \varphi_{h+1} \equiv X_h^2 \pmod{\sigma_h^2 \varphi_h}, \quad (h = 1, 2, \dots, n-1)$$

welche leicht aus einem bekannten Determinantensatze zu erschließen sind, zu Gleichungen:

$$\left(\frac{-\sigma_{h-1} \sigma_h \sigma_{h+1} \varphi_{h-1} \varphi_{h+1}}{\varepsilon_h \varphi_h}\right) = 1; \quad (h = 1, 2, \dots, n-1)$$

die ε_h sollen hier die Vorzeichen der Größen φ_h vorstellen; und das Produkt dieser $n-1$ Gleichungen kann durch wiederholte Anwendung des quadratischen Reziprozitätsgesetzes, wenn man noch zu Hilfe nimmt, daß der Trägheitsindex I mit der Anzahl der Zahlen -1 in der Reihe $\varepsilon_1, \frac{\varepsilon_2}{\varepsilon_1}, \frac{\varepsilon_3}{\varepsilon_2}, \dots, \frac{(-1)^I}{\varepsilon_{n-1}}$ identisch ist, und wenn man mit j die Anzahl der in der Determinante von f in ungeraden Potenzen aufgehenden Primzahlen von der Form $4l+3$ bezeichnet, in die Relation umgewandelt werden:

$$(7) \quad \prod C_p = (-1)^{\left[\frac{n-2I-2j}{4}\right] + \left[\frac{n-2I-2j}{2}\right]},$$

wo das Produkt über die Primzahl 2 und alle in der Determinante von f vorkommenden ungeraden Primzahlen zu erstrecken ist und noch über beliebige weitere Primzahlen ausgedehnt werden kann; die eckige Klammer bedeutet das Symbol für größte ganze Zahlen.

Diese Relation verhilft uns nun zu dem noch fehlenden Nachweise, daß eine Einheit C_p auch bei allen solchen rationalen umkehrbaren Transformationen von f in andere ganzzahlige Formen ungeändert bleibt, bei welchen nicht sowohl Determinante wie Generalnenner zu p relativ prim sind. Eine jede solche Transformation läßt sich nämlich, worauf ich sofort noch näher eingehen werde, darstellen als Produkt aus einer ersten Transformation, in deren Determinante und Generalnenner die Primzahl p ausschließlich eingeht, und aus einer zweiten Transformation, in welcher Determinante und Generalnenner zu p relativ prim sind. Führt die ganze Transformation die Form f wieder in eine ganzzahlige Form über, so

muß solches auch die erste Teiltransformation tun, weil die zweite nicht imstande sein würde, einen durch die erste in f hineingebrachten Generalnenner wieder herauszuschaffen. Die erste Teiltransformation ändert C_p nicht, weil sie nach dem bereits Bewiesenen nur diese Einheit allein ändern könnte, nun aber infolge von (7) eine einzelne solche Einheit sich nicht ändern kann; die zweite ist auf C_p sicher ebenfalls ohne Einfluß, also gilt dasselbe von der ganzen Transformation.

Von der Gleichung für das Produkt aller Einheiten C_p läßt sich noch eine interessante Anwendung machen. Stellt man eine positive Zahl N , in ihre verschiedenen Primzahlpotenzen zerlegt, in der Form $N = \prod p^t$ dar, so ist nach F. Q., p. 52 [[S. 48]]:

$$f(\alpha, N) = \prod f\left(\alpha \frac{N}{p^t}, p^t\right).$$

Bringt man diese Relation mit den Gleichungen (1), (2) und (7) in Verbindung, so folgt:

Für jeden positiven Modul N , welcher durch $\frac{4\Delta}{\sigma_{n-1}d_{n-2}}$ teilbar ist, und für beliebige zu ihm relativ prime Zahlen α ist:

$$\begin{aligned} & \sum e^{\frac{2\pi i \alpha f(x_1, x_2, \dots, x_n)}{N}} \\ &= \left(\frac{(-1)^{\left[\frac{n}{2}\right]} \Delta N^n}{\alpha}\right) i^{-n^2 \left(\frac{\alpha-1}{2}\right)^2 + \left[\frac{n-2I}{2}\right]} \left(\frac{1+i}{\sqrt{2}}\right)^{n-2 \left[\frac{n}{2}\right]} \sqrt{(-1)^I \Delta (2N)^n} \\ & \quad \left(x_h = 1, 2, \dots, N\right) \\ & \quad \left(h = 1, 2, \dots, n\right) \end{aligned}$$

Die erste Klammer auf der rechten Seite bedeutet das Legendresche Symbol. Bemerkenswert ist namentlich, daß diese Summe nur von der Determinante Δ und dem Reste von $I \pmod{4}$ abhängt, im übrigen aber von den Charakteren der Form f völlig unabhängig ist.

Wegen des soeben herangezogenen Hilfssatzes über die Zerlegung beliebiger rationaler Transformationen in solche, deren Determinante und Generalnenner nur durch eine Primzahl aufgehen, kann auf die Aufsätze von Smith „On systems of linear indeterminate equations and congruences“*) und von Frobenius über die „Theorie der linearen Formen mit ganzen Koeffizienten“**) verwiesen werden; doch sind vielleicht auch die folgenden Bemerkungen über jenen Satz nicht uninteressant. Es genügt offenbar, wenn man die in Rede stehenden Zerlegungen für ganz-

*) Philosophical Transactions, vol. 151. 1861. (Collected Papers, vol. I, p. 367.)

**) Crelles Journal, Bd. 86.

zahlige Transformationen auszuführen imstande ist, da man einen etwa vorhandenen Generalnenner jederzeit in angemessener Weise in Faktoren zerlegen und diese dann unter die Teiltransformationen verteilen kann. Jede ganzzahlige Transformation läßt sich nun immer (und zwar auf unendlich viele Arten) darstellen als Produkt aus einer solchen ganzzahligen Transformation, für welche alle Elemente, die in ihrer Determinante rechts von der Diagonale stehen, Null sind und die ich deshalb für einen Moment eine rechts reduzierte Transformation nennen will, und einer ganzzahligen Transformation von der Determinante 1; dieses erhellt aus dem Umstande, daß man eine jede ganzzahlige Determinante auf die in Frage kommende Gestalt in der Weise bringen kann, daß man in ihr wiederholt Vertikalreihen zueinander addiert oder voneinander subtrahiert. Das Produkt zweier rechts reduzierter Transformationen:

$$|p_h^k| \quad (p_h^k = 0, h < k) \quad \text{und} \quad |q_h^k| \quad (q_h^k = 0, h < k) \quad (h, k = 1, 2, \dots, n)$$

ist nun immer wieder eine rechts reduzierte Transformation:

$$|r_h^k| \quad (r_h^k = 0, h < k), \quad (h, k = 1, 2, \dots, n)$$

und zwar ist für diese:

$$r_h^h = p_h^h q_h^h, \quad r_h^{h-d} = \sum p_h^{h-d} q_{h-d}^{h-d} \quad \begin{pmatrix} h=1, 2, \dots, n \\ d=1, 2, \dots, h-1 \\ d=0, 1, \dots, d \end{pmatrix}$$

Umgekehrt ersieht man aus diesen Gleichungen, daß jede rechts reduzierte Transformation in zwei andere zerlegt werden kann, sowie für die Zahlen r_h^h ihrer Diagonalreihe eine Zerlegung in Zahlenpaare p_h^h, q_h^h von solcher Art vorliegt, daß ein jedes p_h^h zu $q_1^1, q_2^2, \dots, q_{h-1}^{h-1}$ relativ prim ist; denn alsdann kann man der Reihe nach für $d = 1, 2, \dots, h-1$ alle Größen p_h^{h-d}, q_h^{h-d} als ganze Zahlen finden, indem man immer zuerst p_h^{h-d} der Kongruenz

$$p_h^{h-d} \equiv \frac{r_h^{h-d} - \sum_{\delta=1}^{d-1} p_h^{h-\delta} q_{h-\delta}^{h-\delta}}{q_{h-d}^{h-d}} \pmod{p_h^h} \quad (h = d+1, d+2, \dots, n)$$

gemäß wählt und hernach

$$q_h^{h-d} = \frac{r_h^{h-d} - \sum_{\delta=1}^d p_h^{h-\delta} q_{h-\delta}^{h-\delta}}{p_h^h} \quad (h = d+1, d+2, \dots, n)$$

setzt. Beispielsweise kann man für die Zahlen p_h^h die höchsten in den Zahlen r_h^h überhaupt aufgehenden Potenzen irgendeiner Primzahl p nehmen, wodurch man auf die für uns wichtige Zerlegung kommt.

Mit dem Nachweis der invarianten Natur der Einheiten C_p ist zugleich die Existenz der oben vermittleis dieser Einheiten definierten Zahl B sichergestellt, und es erübrigt nur noch zu zeigen, daß in der Tat mit den Zahlen n, I, A, B das System derjenigen Funktionen einer quadratischen Form, welche bei allen rationalen umkehrbaren Transformationen der Form ungeändert bleiben, vollständig erschöpft ist.

Zu dem Ende gehe ich wieder von irgendeiner ganzzahligen quadratischen Form f aus, und ich suche dieselbe rational in eine ganzzahlige Form mit möglichst kleiner Determinante zu transformieren. Es fragt sich, kann dabei ein bestimmter Primfaktor p der Determinante in Wegfall kommen oder nicht.

Es sei zunächst p eine ungerade Primzahl. Man bestimme irgendeine mit f äquivalente Grundform in bezug auf p ; aus einer solchen folgt für einen jeden Modul p^t ein Hauptrest von dem unter (5) angegebenen Typus; es werde eine solche Potenz p^t gewählt, welche die zu f gehörige Potenz p^{v_n-1} überschreitet. Sowie dann in dem betreffenden Hauptreste (5) einer der Exponenten v_{h-1} sich ≥ 2 erweisen sollte, kann derselbe um 2 verringert werden, dadurch daß man die zugehörige Variable dem p^{ten} Teile einer neuen Variable gleich setzt, bei welcher Operation alle Koeffizienten des Hauptrestes ganze Zahlen bleiben werden. Indem man diese Reduktion so oft als angänglich wiederholt und am Schlusse nötigenfalls noch eine Umstellung der Variablen vornimmt, kommt man zu einer Form mit einem Reste:

$$\alpha_1 \xi_1^2 + \dots + \alpha_{n-m} \xi_{n-m}^2 + p(\alpha_{n-m+1} \xi_{n-m+1}^2 + \dots + \alpha_n \xi_n^2) \pmod{p^2},$$

worin alle α_h zu p relativ prim sind. Die Determinante dieser Form enthält genau die Potenz p^m als Faktor, und m ist eine Zahl zwischen 0 und n . Ob m gerade oder ungerade ausfällt, wird davon abhängen, ob die Primzahl p in der Zahl A von f nicht enthalten war ($\delta = 0$) oder in ihr vorkam ($\delta = 1$). Die Einheit C_p erhält hier den Ausdruck

$$(8) \quad \left(\frac{(-1)^{\left[\frac{m}{2}\right]} \alpha_{n-m+1} \dots \alpha_n}{p} \right).$$

Im Falle $n = 2$ ist daher, sowie p in A nicht vorkommt und $-A$ quadratischer Rest von p ist, immer $C_p = 1$ (nämlich sowohl für $m = 2$ wie für $m = 0$).

Die erhaltene Form kann nun unter Umständen so transformiert werden, daß an die Stelle von m eine kleinere Zahl tritt. Hat man nämlich eine binäre Form

$$\begin{pmatrix} p\alpha & 0 \\ 0 & p\beta \end{pmatrix} \pmod{p^2}$$

(ich bezeichne hier die Form durch das quadratische Schema ihrer Koeffizienten), in welcher α und β zu p relativ prim sind, und ist $-\alpha\beta$ quadratischer Rest von p , so kann man eine Zahl η finden, so daß $\alpha + \beta\eta^2$ durch p , aber nicht durch p^2 aufgeht, und jene Form verwandelt sich durch die Substitution $\begin{vmatrix} 1 & 0 \\ \eta & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 \\ 0 & p \end{vmatrix}$ in eine andere Form, deren Determinante genau die Potenz p^4 enthält, von der sich aber der quadratische Faktor p^2 ablöst, den man durch die Substitution $\begin{vmatrix} p^{-1} & 0 \\ 0 & p^{-1} \end{vmatrix}$ beseitigen

kann; es resultiert dann eine Form, deren Determinante überhaupt nicht durch p aufgeht und welche wieder eine Grundform in bezug auf p ist.

Ist aber in jener binären Form $-\alpha\beta$ quadratischer Nichtrest von p , so fällt $\alpha + \beta\eta^2 \pmod{p}$ für alle $\frac{p-1}{2}$ modulo p inkongruenten und von Null verschiedenen Quadrate η^2 von Null und von α verschieden aus, und muß deshalb für mindestens eines dieser η einen anderen quadratischen Charakter in bezug auf p liefern als α . Durch die mit dem betreffenden η gebildete Substitution $\begin{vmatrix} 1 & 0 \\ \eta & 1 \end{vmatrix}$ wird dann aus $\begin{pmatrix} p\alpha & 0 \\ 0 & p\beta \end{pmatrix} \pmod{p^2}$ eine Grundform in bezug auf p , die sich in einen analogen Hauptrest überführen läßt, in welchem aber die anstelle von α tretende Zahl einen anderen quadratischen Charakter in bezug auf p hat als α . Aus dieser letzten Bemerkung ersieht man, daß, wenn die Zahl m oben ≥ 3 ist, man es immer so einrichten kann, daß unter den letzten m Zahlen α_h sich zwei solche befinden, deren negativ genommenes Produkt quadratischer Rest von p ist, und welche also zu der vorher beschriebenen Reduktion Gelegenheit bieten. Ob dann, wenn m gerade und bereits auf 2 gebracht ist, jene Operation noch einmal anwendbar erscheint, ob also die Primzahl p durch rationale Transformation aus der Determinante ganz ausgeschafft werden kann oder aber die Determinante immer durch p^2 teilbar bleiben muß, wird davon abhängen, ob $C_p = 1$ ist oder $= -1$. Man sieht demnach, daß, je nachdem die Determinante von f eine gerade Potenz von p enthält und $C_p = 1$ ist, oder sie eine ungerade Potenz von p enthält, oder eine gerade und $C_p = -1$ ist, man von f zu einer Form g gelangen kann, deren Determinante überhaupt nicht durch p aufgeht, oder den Faktor p enthält, oder den Faktor p^2 . Für diese Form g existieren dann außer c_p und C_p keine weiteren quadratischen Charaktere in bezug auf die Primzahl p . Denn alle derartigen Charaktere müßten sich, wie bereits oben bemerkt wurde, aus den Werten der dieser Form zugehörigen Gaußschen Summen $g(\alpha, p^t)$ erschließen lassen; für diese gilt hier aber bereits von $t = 1$ an (in dem ersten der drei unterschiedenen Fälle sogar von $t = 0$ an) die Formel (1).

Mit der so gewonnenen Form können nun entsprechend den weiteren in ihrer Determinante etwa enthaltenen ungeraden Primzahlen analoge Operationen vorgenommen werden, und Potenzen dieser Primzahlen nach Möglichkeit aus der Determinante entfernt werden. (Bis man dabei zur Betrachtung einer bestimmten Primzahl p gekommen ist, hat man lauter Substitutionen angewandt, deren Determinanten zu dieser Primzahl relativ prim sind, und sind deshalb die auf diese Primzahl bezüglichen Potenzen p^{v_h-1} ($h = 1, 2, \dots, n$) von f völlig unberührt geblieben.)

Jetzt betrachte ich das Verhältnis der vorgelegten Form f zur Primzahl 2. Es werde zu f oder zu einer der späteren Formen eine äquivalente Grundform in bezug auf 2 aufgesucht; aus einer solchen entspringt für jeden Modul 2^i ein Hauptrest von dem unter (6) angegebenen Typus, welcher aus Formen mit einer Variable und aus binären Formen $2^v \cdot \begin{pmatrix} 2\alpha & \beta \\ \beta & 2\gamma \end{pmatrix}$ mit ungeraden β und α zusammengesetzt erscheint. Eine jede dieser binären Formen geht durch die Substitution $\begin{vmatrix} 1 & 0 \\ 0 & 2 \end{vmatrix}$ in eine Grundform in bezug auf 2 über, deren zugehörige Hauptreste sich noch weiter in je zwei Formen mit einer Variable auflösen. Man kann so zu einer Grundform in bezug auf 2 gelangen, welche in bezug auf Moduln 2^i Hauptreste ergibt, die in lauter Formen mit einer Variable zerfallen, und von diesen Hauptresten kommt man, ähnlich wie bei ungeraden Primzahlen, zu Formen mit Resten:

$$\alpha_1 \xi_1^2 + \dots + \alpha_{n-m} \xi_{n-m}^2 + 2(\alpha_{n-m+1} \xi_{n-m+1}^2 + \dots + \alpha_n \xi_n^2) \pmod{16},$$

worin alle α_h ungerade sind. Hier erhält die Einheit C_2 den Ausdruck

$$(9) \quad (-1)^{\left[\frac{n}{4}\right] + \left[\frac{n}{2}\right]} \left(\left[\frac{n}{2}\right] + \sum \frac{\alpha_h - 1}{2} \right) + \sum \frac{\alpha_h - 1}{2} \frac{\alpha_h - 1}{2} \left(\frac{2}{\alpha_{n-m+1} \dots \alpha_n} \right) \cdot \begin{pmatrix} h=1, 2, \dots, n \\ k=1, 2, \dots, h-1 \end{pmatrix}$$

Um eine möglichst kleine Zahl m zu erzielen, bemerke ich zunächst, daß eine binäre Form $\begin{pmatrix} 2\alpha & 0 \\ 0 & 2\beta \end{pmatrix} \pmod{8}$, in welcher α und β ungerade

und $\alpha \equiv \beta \pmod{4}$ ist, durch die Substitution $\begin{vmatrix} 1 & 0 \\ 1 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 \\ 0 & 2 \end{vmatrix}$ in eine Form mit dem quadratischen Faktor 4 übergeht; entfernt man diesen

durch die Substitution $\begin{vmatrix} \frac{1}{2} & 0 \\ 0 & \frac{1}{2} \end{vmatrix}$, so bleibt eine Grundform in bezug auf 2

mit ungerader Determinante übrig, die gleichfalls zu zerfallenden Haupt-

resten, aber mit ungeraden Koeffizienten in der Diagonale Veranlassung gibt. Diese Reduktion ist immer möglich, wenn $m \geq 3$ ist, weil dann unter den m letzten Zahlen α_n sich immer mindestens zwei modulo 4 kongruente finden müssen. Ist man so in den Fällen $n \geq 3$ und, wenn m gerade ist, bis zu einem $m = 2$ gelangt, und ist diese Reduktion dann zunächst nicht weiter möglich, so kann man aus dem Hauptreste, den man gerade vor sich hat, gewiß einen Teilrest $\begin{pmatrix} \alpha & 0 \\ 0 & 2\beta \end{pmatrix} \pmod{16}$ mit ungeraden Zahlen α und β herausnehmen, und ein solcher geht durch die Substitution $\begin{vmatrix} 1 & 0 \\ 1 & 1 \end{vmatrix}$ in eine Grundform über, welche Hauptreste ergibt, die anstelle der Zahl β ungerade Zahlen von anderem quadratischen Charakter in bezug auf 4 als β enthalten, wodurch die beschriebene Reduktion noch einmal anwendbar wird. So kann man in den Fällen $n \geq 3$ von f stets zu einer Form g gelangen, deren Determinante entweder ungerade ist oder den Faktor 2 höchstens einmal enthält, und zwar zu einer solchen Form, welche auch ungerade Zahlen darstellt.

Im Falle $n = 2$ würde die hier benutzte Methode zur Verringerung der Zahl m ihren Dienst bei einem Reste $\begin{pmatrix} 2\alpha & 0 \\ 0 & 2\beta \end{pmatrix} \pmod{16}$ versagen, für welchen $-\alpha\beta \equiv 1 \pmod{4}$ ist. Ist dann $-\alpha\beta \equiv 1 \pmod{8}$, so kann man eine Zahl η finden, so daß $\alpha + \beta\eta^2$ durch 8, aber nicht durch 16 aufgeht, und dieser Rest geht durch die Substitution $\begin{vmatrix} 1 & 0 \\ \eta & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 \\ 0 & 8 \end{vmatrix}$ in einen Rest über, von welchem sich der quadratische Faktor 16 rational abtrennen läßt, so daß eine Form g übrig bleibt, die in einen Hauptrest mit einer Zahl $m = 0$ überzuführen ist. *In diesem Falle, also wenn $-\alpha\beta \equiv 1 \pmod{8}$, ist daher immer $C_2 = 1$.*

Hat man aber $-\alpha\beta \equiv 5 \pmod{8}$, so sind die Fälle $m = 0$ und $m = 2$ wesentlich verschieden, und der erste ist mit $C_2 = 1$, der zweite mit $C_2 = -1$ verbunden. Auch der im zweiten Falle sich ergebende Rest $\begin{pmatrix} 2\alpha & 0 \\ 0 & 2\beta \end{pmatrix} \pmod{16}$, $-\alpha\beta \equiv 5 \pmod{8}$ kann in eine Form g mit ungerader Determinante transformiert werden, nämlich durch die Substitution $\begin{vmatrix} 1 & 0 \\ 1 & 1 \end{vmatrix} \cdot \begin{vmatrix} 1 & 0 \\ 0 & 2 \end{vmatrix}$ und nachherige Abtrennung des Faktors 4; diese Form ist dann aber notwendig von der zweiten Art (sie stellt nur gerade Zahlen dar).

Die Formen g , auf welche man so in jedem Falle kommt, deren Determinanten nicht durch 4 aufgehen, besitzen außer den Einheiten c, c_2

und C_2 keine weiteren Charaktere in bezug auf die Primzahl 2, wie aus dem Umstande zu entnehmen ist, daß für sie die Gaußschen Summen $g(\alpha, 2^t)$, soweit dieselben für ein $t > 0$ sich noch nicht der Formel (2) anpassen, einfach Null sind.

Fassen wir alle diese Resultate zusammen, so ist in der Tat gezeigt, daß jede vorgelegte Form f rational in eine Form eines durch ihre Zahlen A und B völlig bestimmten Geschlechts von der zu diesen Zahlen gehörigen Determinante ABB (B bedeutet den Quotienten aus B und dem größten Divisor von A und B) transformiert werden kann, und zwar, mit Ausnahme des Falles $n = 2$, $-A \equiv 5 \pmod{8}$, $C_2 = -1$, in eine Form der ersten Art.

Wenn $n > 2$ ist, kann unter Umständen noch ein zweites Geschlecht von der Determinante ABB und mit denselben Invarianten I, A, B vorhanden sein, nämlich wenn diese Determinante und diese Invarianten ganzzahligen, uneigentlich primitiven Formen angehören können; die Hauptreste dieses Geschlechts in bezug auf Moduln 2^t müßten dann aus lauter

Formen $\begin{pmatrix} 2\alpha & \beta \\ \beta & 2\gamma \end{pmatrix}$ mit ungeraden α und β und eventuell noch einer Form

$2\alpha_5^2$ mit ungeradem α zusammengesetzt sein, was zu den Bedingungen $n \equiv 0$, $A \equiv 1 \pmod{2}$, $c = 1$, $c_2 = C_2$ oder $n \equiv 1$, $A \equiv 0 \pmod{2}$, $c_2 = C_2$ führen würde. Haben diese Beziehungen statt, so existiert jenes zweite Geschlecht wirklich, und wählt man dann, was für $n > 2$ immer möglich ist, einen jener Hauptreste so, daß die Zahl γ in irgendeinem seiner binären Teilreste durch 2, aber nicht durch 4 aufgeht, so kommt man durch eine mit den Variablen dieses Teilrestes vorgenommene Trans-

formation $\begin{vmatrix} 2 & 0 \\ 0 & \frac{1}{2} \end{vmatrix}$ auf eine Form des ersten, eigentlich primitiven Geschlechts zurück.

ZUR GEOMETRIE DER ZAHLEN

VIII.

Über die positiven quadratischen Formen und über kettenbruchähnliche Algorithmen.

(Crelles Journal für die reine und angewandte Mathematik, Band 107, S. 278—297.)

Die wesentlich positiven quadratischen Formen verdienen und gestatten eine besondere Behandlung durch den Umstand, daß sie die einfachsten Formen sind, bei welchen durch den Wert der Form zugleich die Werte sämtlicher Veränderlichen begrenzt sind. Aus diesem Grunde erscheinen sie als ein naturgemäßes Hilfsmittel für die Untersuchung von Reihen diskreter Größen, und in diesem Sinne sind sie namentlich von Herrn Hermite zu wiederholten Malen mit bedeutendem Erfolge verwendet worden.

Wenn ihren Koeffizienten auch ganz beliebige reelle Werte, nicht durchaus rationale beigelegt werden, so stellen sie doch immer geeignete Formen für Zahlen vor, d. h. es hat einen Sinn, die Unbestimmten in ihnen auf die Reihe der ganzen Zahlen zu beschränken. Bei einer solchen Auffassung können diese Formen im speziellen als der analytische Ausdruck gewisser einfacher geometrischer Gebilde gelten, der parallelepipedisch angeordneten regelmäßigen Punktsysteme, und es müssen irgend zwei Formen als äquivalent betrachtet werden, welche auseinander durch lineare Substitutionen mit ganzzahligen Koeffizienten und von einer Determinante ± 1 hervorgehen.

Nun entsteht die Aufgabe, eine Vereinigung äquivalenter Formen, eine Klasse, vollständig durch Invarianten zu charakterisieren. Erst für binäre Formen hat durch die Untersuchungen von Herrn Kronecker diese Aufgabe insofern eine vollkommene Lösung gefunden, als hier in hinreichender Anzahl invariante Bildungen als explizite Funktionen eines beliebigen Elements der Klasse und in einer Gestalt, welche die Invarianteneigenschaft unmittelbar in Evidenz treten läßt, gewonnen sind. Ähnliche Aufschlüsse hinsichtlich der Formen mit höherer Variablenzahl mögen aus den jüngsten Arbeiten dieses Forschers zu erwarten sein.

Indes ist die genannte Aufgabe einer Lösung noch in einem anderen

Sinne fähig. Gelingt es, aus den unendlich vielen Formen einer Klasse durch bestimmte Bedingungen eine einzige auszusondern, so stellt eine solche sogenannte reduzierte Form gewissermaßen ebenfalls ein vollständiges Invariantensystem der Klasse vor, nur daß der Ausdruck dieses Systems von irgendeiner gegebenen Form der Klasse auch jedesmal erst durch ein gewisses besonderes Verfahren (das dafür aber nur arithmetische Operationen in beschränkter Zahl erfordern darf) hergeleitet werden kann.

In solcher Art hat Lagrange*) die Theorie der binären quadratischen Formen in Angriff genommen und zu einem glänzenden Abschlusse gebracht. Seine Resultate über die definiten Formen erhielten durch Legendre**) eine Fassung, welche wohl auf ihre Verallgemeinerungsfähigkeit hinweisen konnte.

Aus der fünften Sektion der „Disquisitiones arithmeticae“ entnahm Seeber***) die Anregung zu einem Studium der analogen Fragen betreffs der ternären definiten Formen. Seine äußerst mühsame und nicht erfolglose Arbeit fand eine angemessene Würdigung in einer von Gauß†) herrührenden, höchst bemerkenswerten Anzeige. Namentlich durch zweierlei ist diese Anzeige ausgezeichnet: einmal durch den Hinweis auf das von uns schon erwähnte geometrische Äquivalent einer Klasse von positiven quadratischen Formen, dann durch eine eigentümliche Identität, mittels deren eine wichtige, von Seeber nur durch Induktion gefundene Grenze für die Koeffizienten seiner reduzierten Formen direkt in Erscheinung tritt.

Die beschwerliche Methode und die verwickelten Beweise von Seeber veranlaßten Dirichlet††), für den das nicht Einfache überall nur ein Zeichen des Unvollkommenen war, zu einer von Grund aus neuen Behandlung, bei welcher er besonders auch durch die von Gauß nur mehr in ihren Umrissen angedeutete geometrische Einkleidung eine außerordentliche Durchsichtigkeit erzielte. Der große Fortschritt von Dirichlet bestand darin, daß er nicht mit dem schwerfälligen rechnerischen Ausdrucke der Ungleichungen operierte, durch welche Seeber reduzierte Formen definiert hatte, sondern mit deren wohlerkannter innerer Bedeutung, welche darauf hinausging, die reduzierte Form von gewissen in dem zu-

*) Recherches d'Arithmétique. Mémoires de l'Académie de Berlin, 1773, p. 265. (Oeuvres, T. III, p. 695.)

**) Théorie des Nombres, 3^{me} éd., T. I § VIII.

***) Untersuchungen über die Eigenschaften der positiven ternären quadratischen Formen, Freiburg i. B., 1831.

†) Göttingische gelehrte Anzeigen, Jahrg. 1831, 2. S. 1065 (auch Crelles Journal, Bd. 20, S. 312 und Werke, Bd. II, S. 188).

††) Ueber die Reduction der positiven quadratischen Formen mit drei unbestimmten ganzen Zahlen, Crelles Journal, Bd. 40, 1850, S. 209. (Werke, Bd. II, S. 21.)

gehörigen Punktsysteme vorkommenden kleinsten Entfernungen abhängig zu machen.

Dasselbe ebenso einfache wie sachgemäße Prinzip, doch in rein arithmetischer Fassung, befolgte Herr Hermite*) in seinen zahlentheoretischen Briefen an Jacobi, welche in demselben Bande von Crelles Journal gedruckt sind, in dem die ausführlichere Darstellung Dirichlets nach einem bereits vorher im Monatsbericht der Akademie (Jahrg. 1848) gegebenen Auszuge erschien. Die Untersuchungen von Herrn Hermite beziehen sich auf Formen mit beliebiger Variablenzahl; sie beginnen mit der Aufstellung des Fundamentalsatzes der Reduktion, wonach die kleinste, durch eine positive quadratische Form von n Variablen mittels ganzer Zahlen darstellbare, von Null verschiedene Größe in ihrem dimensionslosen Verhältnis zur n^{ten} Wurzel aus der Determinante der Form niemals einen gewissen, nur von der Zahl n abhängigen Betrag übersteigt; und sie stellen sich in ihrem Verlaufe als ein ununterbrochenes Zeugnis für die Fruchtbarkeit dieses Satzes in fast jedem Abschnitte der Zahlenlehre dar; es seien nur die Anwendungen auf Kettenbrüche, komplexe Einheiten und approximative Auflösung von Gleichungen hervorgehoben.

Insbesondere ergibt sich aus jenem Satze mit Leichtigkeit und noch auf mannigfache Weise die Endlichkeit der Klassenanzahl bei Beschränkung auf ganzzahlige Werte der Koeffizienten und einen festen ganzzahligen Wert der Determinante. Für diese spezielle Folgerung mußte offenbar bereits ein Verfahren genügen, um aus jeder Klasse überhaupt nur eine endliche Anzahl von Formen, nicht gerade eine einzige auszuwählen. Eine wertvolle Ergänzung lieferte deshalb Herr Camille Jordan**) durch den Nachweis, daß bei gewissen Festsetzungen wenigstens eine bloß von der Variablenzahl abhängige Grenze für die Anzahl der im Maximum aus einer Klasse ausgesonderten Formen besteht, indem überhaupt die Substitutionen, durch welche die ausgesonderten Formen ineinander bei Äquivalenz oder in sich selbst übergehen könnten, von vornherein mit der Variablenzahl und zwar in beschränkter Anzahl angewiesen erscheinen.

Neue Gesichtspunkte eröffneten Korkine und Zolotareff***), indem sie jene besonderen Formen heranzogen und bis zur Variablenzahl fünf vollständig bestimmten, für welche das in dem Fundamentalsatze von Hermite genannte Verhältnis (des durch ganze Zahlen erreichbaren Minimum zur n^{ten} Wurzel aus der Determinante) ein Maximum ist.

*) Crelles Journal, Bd. 40, 1850, S. 261—315. (Oeuvres, T. I, pp. 100—163.)

**) Mémoire sur l'équivalence des formes. Journal de l'École Polytechnique T. XXIX, Cah. 48, 1880, p. 111.

***) Mathematische Annalen, Bd. 6, 1873, S. 366 und Bd. 11, 1877, S. 242.

In dem vorliegenden Aufsätze versuche ich hauptsächlich, gewisse Lücken auszufüllen, welche sich in der Theorie der positiven quadratischen Formen gegenwärtig noch fühlbar machen. So geht bei den bisher eingeführten reduzierten Formen mit höheren Variablenanzahlen der den ursprünglichen binären reduzierten Formen von Lagrange innewohnende Charakter verloren, durch eine Reihe von linearen Ungleichungen in den Koeffizienten definiert zu sein. Es erscheint mir aber theoretisch als eine Tatsache von ganz hervorragender Bedeutung, daß man instande ist, aus der $\frac{1}{2}n(n+1)$ -fachen Mannigfaltigkeit, in welcher eine jede quadratische Form von n Variablen durch einen Punkt, unter Zugrundelegung der Werte der Koeffizienten als Koordinaten, repräsentiert wird, aus dieser Mannigfaltigkeit mit Hilfe einer beschränkten Anzahl von lauter ebenen $(\frac{1}{2}n(n+1) - 1)$ -fachen Mannigfaltigkeiten ein zusammenhängendes Gebiet abzugrenzen, in welchem — die Grenzen sind nur teilweise mit einzurechnen — jeder Punkt je eine Klasse von positiven Formen vertritt, und jede solche Klasse auch einmal und nur einmal vertreten ist.

Ein solches Gebiet wird durch die $(\frac{1}{2}n(n+1) - 1)$ -fache Mannigfaltigkeit aller Formen von einer festen positiven, im übrigen beliebigen Determinante in zwei Teile geschieden, von denen der am Nullpunkt befindliche einen endlichen Inhalt hat. Der Ausdruck dieses Inhalts wird hier allgemein mitgeteilt. Es steht dieser Inhalt mit interessanten mittleren Werten der Zahlentheorie im Zusammenhang.

Die Überführung irgendeiner gegebenen Form in eine reduzierte muß durch ausschließliche Verwendung einer beschränkten Zahl a priori anzuweisender Operationen geleistet werden können, und die Ausgangsform darf jedesmal nur in bezug auf Reihenfolge und Wiederholung der Operationen maßgebend sein; dieser berechtigten Forderung wird hier genügt werden.

Mit Hilfe einer auch auf Formen mit mehr als drei Veränderlichen übertragenen geometrischen Ausdrucksweise gelingt es, den Fundamentalsatz von Hermite über das Minimum einer positiven quadratischen Form nicht allein als in gewissem Sinne evident hinzustellen, sondern auch die in diesem Satze und in Erweiterungen desselben benötigten Grenzen den bisher angezeigten gegenüber beträchtlich zu verengern. Dadurch werden dann neue, folgenreiche Anwendungen dieses Satzes möglich. —

1. Die Grundeigenschaft der wesentlich positiven quadratischen Formen.

Eine wesentlich positive quadratische Form kann nur für eine endliche Anzahl von ganzzahligen Wertsystemen ihrer Veränderlichen Werte annehmen, die eine gegebene GröÙe nicht überschreiten.

Denn eine solche Form f mit n Variablen $x_1, x_2, \dots, x_n$ ist bekanntlich immer als Summe der Quadrate von n unabhängigen reellen Linearformen ihrer Variablen darstellbar:

$$f = \xi_1^2 + \xi_2^2 + \dots + \xi_n^2,$$

$$(1a) \quad \xi_a = \pi_{a1}x_1 + \pi_{a2}x_2 + \dots + \pi_{an}x_n, \quad |\pi_{ab}| \neq 0. \quad (a, b = 1, 2, \dots, n)$$

Eine Ungleichung $f \leq G$ mit positivem G hat nun für jede der Linearformen abs. $\xi_a \leq \sqrt{G}$ zur Folge. Lautet die Auflösung dieser Formen nach ihren Variablen:

$$(1b) \quad x_b = \varphi_{1b}\xi_1 + \varphi_{2b}\xi_2 + \dots + \varphi_{nb}\xi_n, \quad (b = 1, 2, \dots, n)$$

so muß daher

$$\text{abs. } x_b \leq \sqrt{G} (\text{abs. } \varphi_{1b} + \text{abs. } \varphi_{2b} + \dots + \text{abs. } \varphi_{nb}) \quad (b = 1, 2, \dots, n)$$

sein. Diesem Systeme von Ungleichungen können aber nur eine endliche Anzahl von ganzzahligen Systemen $x_1, x_2, \dots, x_n$ entsprechen. Natürlich brauchen diese dann nicht sämtlich ein $f(x_1, x_2, \dots, x_n) \leq G$ zu ergeben.

2. Die geometrische Interpretation der wesentlich positiven quadratischen Formen.

In einer ebenen n -fachen Mannigfaltigkeit mögen die Werte der Linearformen $\xi_1, \xi_2, \dots, \xi_n$ solche Koordinaten für einen veränderlichen Punkt P abgeben, daß das Quadrat des von P ausgehenden Linearelements durch die Summe der Quadrate der Differentiale $d\xi_1, d\xi_2, \dots, d\xi_n$ ausgedrückt erscheint. Der Nullpunkt der so vorausgesetzten rechtwinkligen Koordinaten heiße O . Jedem Wertsysteme der Variablen $x_1, x_2, \dots, x_n$ entspricht gemäß (1a) ein Wertsystem $\xi_1, \xi_2, \dots, \xi_n$ und kommt jetzt ein Punkt P zu; die Form $f(x_1, x_2, \dots, x_n)$ stellt offenbar das Quadrat der Entfernung dieses Punktes P von dem festen Nullpunkte O dar.

Welche Punkte gehören nun den ganzzahligen Systemen $x_1, x_2, \dots, x_n$ an? Um diese Punkte zu finden, hat man zuvörderst diejenigen n Punkte $P_1, P_2, \dots, P_n$ kenntlich zu machen, für welche jedesmal eine der n Größen $x_1, x_2, \dots, x_n$ den Wert 1 und die anderen $n - 1$ den Wert Null haben. Liegen diese n Punkte vom Nullpunkte O beziehlich um die Strecken $p_1, p_2, \dots, p_n$ ab, so wird der Punkt, welcher einem willkürlich gewählten Systeme $x_1, x_2, \dots, x_n$ angehört, gefunden, indem man vom Nullpunkte aus die Strecke

$$p_1x_1 + p_2x_2 + \dots + p_nx_n$$

konstruiert, wobei die Additionen in geometrischem Sinne zu verstehen sind.

Hiernach würde man folgendermaßen zu den sämtlichen Punkten mit ganzzahligen Bestimmungsstücken $x_1, x_2, \dots, x_n$ gelangen können: Man stelle in die am Punkte O von den dort ausgehenden n Strecken $p_1, p_2, \dots, p_n$

gebildete n -kantige Ecke — das Vorhandensein einer solchen wirklichen Ecke ist die Folge der Unabhängigkeit der Gleichungen (1a) — ein mit diesen n Strecken als Kanten konstruiertes n -dimensionales Parallelepipedum. Von den $2n(n-1)$ -dimensionalen Begrenzungsebenen dieses Parallelepipedum wollen wir die n durch den Punkt O nicht hindurchgehenden in ihrer ganzen Ausdehnung, d. h. mit allen ihren Grenzlinien verschiedener Dimensionen, also im besonderen mit allen übrigen Eckpunkten, als nicht mehr zu dem Bereiche des Parallelepipedum gehörig betrachten; in ähnlichem Sinne wollen wir uns auch künftighin den Bereich jedes irgend einmal vorkommenden Parallelepipedum festgesetzt denken. An jede der $2n$ Begrenzungsflächen dieses Grundparallelepipedum lege man gleichgerichtet ein vollkommen gleiches Parallelepipedum, an die noch freien Begrenzungsflächen dieser Parallelepipeda wieder ein gleiches, und dieses Verfahren denke man sich unbegrenzt fortgesetzt. Dann finden sich die gesuchten Punkte in den einzelnen Hauptecken dieser nacheinander konstruierten Parallelepipeda.

Das vollständige System dieser Punkte mit ganzzahligen Bestimmungsstücken $x_1, x_2, \dots, x_n$ ist um jeden einzelnen seiner Punkte in gleicher Weise gelagert. Wir nennen es deshalb ein *regelmäßiges Punktsystem*. Wir werden ein solches System mitunter einfach mit dem Buchstaben $\mathfrak{P}$ ohne weiteren Zusatz, oder wenn die Dimension des Systems kenntlich gemacht werden soll, mit $\mathfrak{P}^{(n)}$ bezeichnen. Ein System $\mathfrak{P}$ besetzt nach irgendeiner Parallelverschiebung entweder vollständig neue Punkte oder tritt wieder ganz in die anfänglichen Lagen seiner Punkte ein.

Da die zu konstruierenden Parallelepipeda den ganzen vorausgesetzten n -dimensionalen Raum lückenlos erfüllen werden, und da sie überdies nach den Punkten des Systems zählbar, d. h. ihnen eindeutig zugeordnet sind — nach unseren Festsetzungen über den Bereich dieser Parallelepipeda ist ein jeder Punkt des Raumes einem und nur einem der Parallelepipeda zuzuteilen —, so wird innerhalb eines, überallhin gleichmäßig ins Unendliche ausgedehnten Gebiets (man denke beispielsweise an einen n -dimensionalen Würfel mit unendlich großer Kante) im Durchschnitt ein Punkt des Systems auf einen Raumteil gleich dem Volumen des Grundparallelepipedum kommen. In der Maßzahl dieses Volumens erkennen wir hiernach eine für das Punktsystem an sich charakteristische und von der Wahl des Gerüsts, durch welches wir die Punkte verbunden haben, völlig unabhängige Konstante; und den reziproken Wert dieser Maßzahl werden wir passend als die *mittlere Dichtigkeit des Punktsystems* bezeichnen können.

Zu jedem Punktsysteme gibt es offenbar ein geometrisch ähnliches Punktsystem von der mittleren Dichtigkeit 1.

3. Erneuter Beweis der Grundeigenschaft.

Die in 1. bewiesene Grundeigenschaft der wesentlich positiven quadratischen Formen läßt sich nun auch leicht geometrisch einsehen.

Die gesamten Begrenzungsflächen der vorhin konstruiert gedachten Parallelepipeda sind enthalten in n verschiedenen Scharen von lauter parallelen und äquidistanten $(n-1)$ -dimensionalen Ebenen, als deren Durchschnitte eben die Punkte unseres Systems sich ergeben. Die Distanzen in den einzelnen Scharen werden durch die n Höhen des Grundparallelepipedum geliefert; die Längen dieser Höhen mögen $h_1, h_2, \dots, h_n$ heißen.

In jeder einzelnen Schar sind die Elemente nach einer bestimmten der n Zahlen $x_1, x_2, \dots, x_n$ zu numerieren. Im Nullpunkte O kreuzen sich die Nullelemente aller Scharen; in einem Punkte P mit ganzzahligen Bestimmungsstücken $x_1, x_2, \dots, x_n$ das x_1^{te} Element der ersten, das x_2^{te} der zweiten, $\dots$, das x_n^{te} der n^{ten} Schar. Nun kann der Abstand OP nicht kleiner sein als der senkrechte Abstand zweier durch O und P gehenden $(n-1)$ -dimensionalen Parallelebenen. Soll also der Abstand OP eine gegebene Länge $\sqrt{G}$ nicht überschreiten, d. h. soll:

$$f(x_1, x_2, \dots, x_n) \leq G$$

sein, so müssen um so mehr die Ungleichungen statthaben:

$$(3) \quad h_a \text{ abs. } x_a \leq \sqrt{G}, \quad (a = 1, 2, \dots, n)$$

und diesen kann wieder nur eine beschränkte Anzahl von ganzen Zahlen genügen.

4. Positive quadratische Form und Parallelelepipedum.

Die mit Ausnahme des Falles $n=1$ immer vorhandene Willkür in der Darstellung einer positiven quadratischen Form f als Summe von n Quadraten linearer Formen betrifft geometrisch nur die Neigung der Elementarparallelepipeda gegen die rechtwinkligen Koordinatenachsen, auf welchen die linearen Formen ihre Auslegung finden: es sind nämlich die Projektionen der Strecken p_b auf die Achsen der $\xi_1, \xi_2, \dots, \xi_n$ genau die Koeffizienten $\pi_{1b}, \pi_{2b}, \dots, \pi_{nb}$ der zugehörigen Variablen x_b in den linearen Formen $\xi_1, \xi_2, \dots, \xi_n$.

Die Figur des Elementarparallelepipedum, ohne Rücksicht auf ihre Stellung im vorausgesetzten n -dimensionalen Raume, aber mit Kennzeichnung ihrer Ursprungscke und der Reihenfolge der Kanten an dieser Ecke, bestimmt eindeutig den Ausdruck der Form f in ihren Koeffizienten. Soll dieser:

$$(4) \quad f = \sum q_{ab} x_a x_b \quad \left(\begin{array}{l} a, b = 1, 2, \dots, n \\ q_{ab} = q_{ba}, a \neq b \end{array} \right)$$

lauten, so bedeutet jedesmal ein Koeffizient q_{aa} mit gleichen Indizes das Quadrat der Länge der Strecke p_a , und ein Koeffizient q_{ab} mit verschiedenen Indizes das Produkt aus den Längen der Strecken p_a und p_b und dem Kosinus des Neigungswinkels dieser Strecken. Ferner bedeuten: 1. die Determinante der Form, $|q_{ab}| = \Delta$, das Quadrat des Inhalts des Parallelepipedum, 2. die symmetrischen Unterdeterminanten $\frac{\partial \Delta}{\partial q_{aa}}$ die Quadrate der Inhalte seiner paarweise einander gleichen Begrenzungsflächen, so daß für die n Höhen des Parallelepipedum die Ausdrücke resultieren:

$$h_a = \sqrt{\frac{\Delta}{\frac{\partial \Delta}{\partial q_{aa}}}}. \quad (a = 1, 2, \dots, n)$$

Alle diese Beziehungen sind am einfachsten durch ein Zurückgehen auf das rechtwinklige Koordinatensystem der $\xi_1, \xi_2, \dots, \xi_n$ einzusehen, werden übrigens sofort noch klarer hervortreten.

Umgekehrt gehören dagegen zur gegebenen wesentlich positiven Form f (Formel (4)) in dem gleichen Raume von n Dimensionen immer zwei verschiedene Arten von n -kantigen begrenzten Ecken, und dementsprechende Parallelepipeda. Denn zunächst haben wir jedenfalls in 2. eine solche Art gefunden, und zwar auf Grund irgendeiner Darstellung von f als Summe der Quadrate von n linearen Formen. Um das dabei angewandte Verfahren beschreiben zu können, ohne auf die Bedeutung der ξ -Koordinaten wieder eingehen zu müssen, wollen wir uns auf den positiven Seiten der rechtwinkligen Koordinatenachsen der $\xi_1, \xi_2, \dots, \xi_n$ die n Punkte $E_1, E_2, \dots, E_n$ markiert denken, welche in der Einheit der Entfernung von O abliegen, und die geometrischen Strecken nach diesen Punkten mit $e_1, e_2, \dots, e_n$ bezeichnen. Dann entsteht die erste, zur Form f gehörige Ecke $O(P_1 P_2 \dots P_n)$ aus 2. einfach, indem in O die Strecken:

$$(4a) \quad p_b = \pi_{1b} e_1 + \pi_{2b} e_2 + \dots + \pi_{nb} e_n \quad (b = 1, 2, \dots, n)$$

angefügt werden. Der günstige Erfolg dieser Operation läßt sich am besten mit Hilfe einer von Graßmann eingeführten Symbolik übersehen: Das Produkt aus den Längen zweier Strecken l und m und dem Kosinus ihres Neigungswinkels mag das *innere Produkt* dieser Strecken heißen und $l|m$ geschrieben werden. Offenbar gilt für diese Art von Multiplikation neben den Regeln $e_a|e_a = 1$, $e_a|e_b = 0$ ($a \neq b$) das distributive Gesetz, und daraus geht sofort $p_a|p_b = q_{ab}$ hervor. Andererseits aber folgt allein aus den Beziehungen $p_a|p_b = q_{ab}$, sowie man mit Hilfe der aus (1b) entnommenen Koeffizienten n Strecken:

$$(4b) \quad e_a = \varphi_{a1} p_1 + \varphi_{a2} p_2 + \dots + \varphi_{an} p_n \quad (a = 1, 2, \dots, n)$$

bildet, mit Notwendigkeit: $e_a|e_a = 1$, $e_a|e_b = 0$ ($a \neq b$), und also jedesmal eine rechtwinklige Ecke mit n Kanten gleich der Längeneinheit.

Von solchen rechtwinkligen Ecken gibt es nun in einem Raume von n Dimensionen (genau so wie im speziellen Falle $n = 3$) immer zwei Arten, die innerhalb dieses Raumes nicht in allen ihren entsprechenden Kanten zugleich zur Deckung zu bringen sind. Eine Art können wir als durch die Ecke $O(E_1 E_2 \dots E_n)$ der Koordinatenachsen definiert betrachten. Dann entsteht die zweite Art durch Spiegelung dieser Ecke an irgend-einer $(n - 1)$ -dimensionalen Ebene, und durch die nämliche Spiegelung muß aus der zuerst gefundenen Ecke $O(P_1 P_2 \dots P_n)$ eine zweite zu f gehörige Ecke hervorgehen. Der Spiegelung an einer Ebene:

$$\varphi_1 \xi_1 + \varphi_2 \xi_2 + \dots + \varphi_n \xi_n = 0$$

entspricht die Umwandlung der Koeffizienten π_{ab} in:

$$\pi_{ab} - 2 \left(\frac{\pi_{1b} \varphi_1 + \pi_{2b} \varphi_2 + \dots + \pi_{nb} \varphi_n}{\varphi_1^2 + \varphi_2^2 + \dots + \varphi_n^2} \right) \varphi_a. \quad (a, b = 1, 2, \dots, n)$$

Man überzeugt sich leicht, daß dabei die Quadratsumme der linearen Formen $\xi_1, \xi_2, \dots, \xi_n$ den Ausdruck f ungeändert beibehält, die Determinante $|\pi_{ab}|$ hingegen in den entgegengesetzten Wert übergeht, was eben das Zeichen dafür ist, daß in der Tat eine Ecke neuer Art gewonnen ist.

Überhaupt können wir nämlich in n Dimensionen zwei Arten von n -kantigen wirklichen Ecken unterscheiden, in diesem Sinne: n -kantige Ecken gleicher Art sind durch kontinuierliche Abänderungen, ohne daß sie aufhören, wirkliche Ecken zu bleiben, zum Zusammenfallen in allen ihren entsprechenden Kanten zu bringen; bei n -kantigen Ecken ungleicher Art ist solches nicht möglich. Daß n Strecken $p_1, p_2, \dots, p_n$ eine *wirkliche* Ecke bilden, heißt soviel, wie daß sie auf ein Parallelepipedum von nicht-verschwindendem n -dimensionalem Inhalte führen. Den fraglichen Inhalt können wir als eine Art von Produkt auffassen und $p_1 p_2 \dots p_n$ schreiben. Wir haben es alsdann mit der sogenannten *äußeren Multiplikation* von Strecken zu tun (es ist das wieder eine Bezeichnung von Graßmann), für welche neben dem distributiven und dem assoziativen Gesetze offenbar die Regel gilt, daß das Produkt einer Strecke in sich selbst Null ist, woraus für zwei Strecken l und m durch Betrachtung von $(l + m)(l + m)$ sich $lm = -ml$ ergibt; und mit Hilfe dieser Regeln folgt aus (4a): $p_1 p_2 \dots p_n = |\pi_{ab}| e_1 e_2 \dots e_n$. Das Nichtverschwinden der Determinante $|\pi_{ab}|$ ist hiernach für die Bildung einer wirklichen Ecke charakteristisch, und nach dem Vorzeichen dieser Determinante richtet sich dann, wie man leicht einsieht, die Art der Ecke.

In dem aus einer Ecke $(p_1, p_2, \dots, p_n)$ folgenden Parallelepipedum befindet sich der betreffenden Ecke diametral gegenüber eine Ecke mit den Kanten $-p_1, -p_2, \dots, -p_n$; diese zweite Ecke gehört offenbar zu derselben quadratischen Form, ergibt aber bei ungeradem n ein entgegen-

gesetztes Kantenprodukt. Bei ungeradem n sind demnach die Parallelepipeda, welche aus zwei zusammengehörigen Ecken ungleicher Art entstehen, im wesentlichen identisch, sie erscheinen nur verschieden aufgefaßt, während bei geradem n eine Deckung solcher zweier Parallelepipeda in der Regel erst innerhalb eines dimensionsreicheren Raumes zu erzielen sein wird. —

Da die quadratische Form f als Ausdruck eines parallelepipedisch geordneten, regelmäßigen Punktsystems in einem gegebenen Raume, wie wir sehen, eine Zweideutigkeit bestehen läßt, so erscheint es vielleicht angebrachter, als solchen Ausdruck die *lineare* Form:

$$p_1 x_1 + p_2 x_2 + \cdots + p_n x_n = p$$

zu nehmen. In dieser sind die Koeffizienten allerdings nicht reine Zahlen, sie bedeuten vielmehr Strecken, bestimmt in Richtung und Länge; aber durch ausschließliche Betrachtung dieser Strecken sind keine weiteren Zahlgrößen zu entnehmen, als eben die Koeffizienten von f . Das Volumen $p_1 p_2 \dots p_n$ setzen wir immer als von Null verschieden voraus. Nach der vorhin erklärten Ausdrucksweise würde f als das innere Produkt dieser linearen Form p in sich selbst, als ihr inneres Quadrat, zu bezeichnen sein.

5. Anschauliche Auslegung des Äquivalenzbegriffs.

Ist ein, auf irgendeine Weise in parallelepipedischer Anordnung gegebenes regelmäßiges Punktsystem $\mathfrak{P}^{(n)}$ noch weiterer solcher Anordnungen fähig?

Bei jeder solchen Anordnung müßte jeder Punkt des Systems als Ausgangspunkt einer, in n anderen Punkten des Systems endenden Ecke eines jedesmal gleichen Parallelepipedium erscheinen, in dessen Bereich — wenn die dem Punkte nicht anliegenden Begrenzungsflächen immer vollständig ausgeschlossen werden — kein weiterer Punkt des Systems fallen dürfte. Verbinden wir also zunächst einen beliebigen Punkt O des Systems mit n beliebigen anderen Punkten des Systems, die nur nicht sämtlich mit O zusammen bereits in einer $(n-1)$ -dimensionalen Ebene liegen sollen. Die Strecken von O nach diesen n Punkten mögen $q_1, q_2, \dots, q_n$ heißen. Da diese Strecken lauter Punkte des Systems verbinden, so werden sie mit den für die gegebene Anordnung charakteristischen Strecken $p_1, p_2, \dots, p_n$ durch irgendwelche Relationen:

$$q_b = s_{1b} p_1 + s_{2b} p_2 + \cdots + s_{nb} p_n \quad (b = 1, 2, \dots, n)$$

mit lauter ganzzahligen Koeffizienten s_{ab} verbunden sein, die nur ein nichtverschwindendes Volumen $q_1 q_2 \dots q_n$ (d. i. eine nichtverschwindende Determinante $|s_{ab}|$) zu ergeben haben; und deshalb wird dann weiter auch

jede beliebige, von O aus konstruierte Strecke $q_1 y_1 + q_2 y_2 + \dots + q_n y_n = q$ mit ganzzahligen Bestimmungsstücken $y_1, y_2, \dots, y_n$ auf einen Punkt des Systems auslaufen müssen.

Ein von O aus, der eben genannten linearen Form q gemäß, in parallelepipedischer Anordnung aufgebautes Punktsystem Ω wird also ganz in dem vorausgesetzten Punktsysteme $\mathfrak{P}$ enthalten sein. Dieses System Ω wird mit $\mathfrak{P}$ zusammenfallen, also eine neue parallelepipedische Anordnung von $\mathfrak{P}$ darbieten, wenn die Parallelepipeda von Ω in ihren Bereichen — dieselben in dem früher festgesetzten Sinne genommen — außer ihrer jedesmaligen Hauptecke keine weiteren Punkte von $\mathfrak{P}$ enthalten. Jedenfalls enthält nun jedes dieser Parallelepipeda gleich viele Punkte aus $\mathfrak{P}$, sagen wir s , und an lauter entsprechenden Stellen, da die Parallelverschiebungen, durch welche Ω mit sich selbst zur Deckung kommt, ja nichts weiter als ein Teil der Deckbewegungen von $\mathfrak{P}$ sind. In 2. sahen wir, daß innerhalb eines unendlich großen n -dimensionalen Würfels aus dem Punktsysteme $\mathfrak{P}$ im Durchschnitt ein Punkt auf einen Raumteil gleich dem Volumen $p_1 p_2 \dots p_n$ kommt, und nun sollen offenbar s solcher Punkte im Durchschnitt auf einen Raumteil gleich dem Volumen $q_1 q_2 \dots q_n = |s_{ab}| p_1 p_2 \dots p_n$ kommen. Mithin kann die Zahl s nur den absoluten Wert der Determinante $|s_{ab}|$ vorstellen, und die Bedingung für eine neue parallelepipedische Anordnung des Punktsystems $\mathfrak{P}$ lautet: $|s_{ab}| = \pm 1$. Es ist geometrisch evident, daß bei Erfüllung dieser Bedingung umgekehrt auch $p_1, p_2, \dots, p_n$ als lineare Funktionen von $q_1, q_2, \dots, q_n$ lauter ganzzahlige Koeffizienten werden aufweisen müssen.

Die Form q geht aus der Form $p = p_1 x_1 + p_2 x_2 + \dots + p_n x_n$ vermittels der Substitution:

$$x_a = s_{a1} y_1 + s_{a2} y_2 + \dots + s_{an} y_n \quad (a = 1, 2, \dots, n)$$

hervor, und es ist klar, daß durch dieselbe Substitution aus der quadratischen Form $p|p = f$ eine quadratische Form g entsteht, welche als das innere Quadrat von q erscheinen wird. Die Eigenschaften der zugehörigen Punktsysteme machen es verständlich, daß man die, durch eine solche lineare Substitution mit ganzzahligen Koeffizienten aus einer Form f hervorgehende Form g als *enthalten* in der Form f bezeichnet, ebenso, daß man sie der Form f *äquivalent* nennt, wenn die Determinante $|s_{ab}| = \pm 1$ ist. Man spricht von *eigentlicher* oder *uneigentlicher Äquivalenz*, je nachdem $|s_{ab}| = 1$ oder $= -1$ ist, je nachdem also die für f und g übereinstimmenden Punktsysteme aus Ecken gleicher oder ungleicher Art herzuleiten sind. Da bei ungeradem n vermöge der Substitution $x_1 = -y_1, x_2 = -y_2, \dots, x_n = -y_n$ jede Form sich selbst auch uneigentlich äquivalent ist, so hat diese letztere Unterscheidung nur bei geradem n einen Wert;

natürlich können auch hier unter Umständen Formen einander eigentlich und uneigentlich äquivalent zu gleicher Zeit sein.

Eine Klasse von äquivalenten Formen entspricht nun dem Inbegriff aller möglichen parallelepipedischen Anordnungen eines Punktsystems $\mathfrak{P}$.

6. Von dem Minimum einer wesentlich positiven quadratischen Form.

In einem parallelepipedisch geordneten, regelmäßigen Punktsysteme $\mathfrak{P}^{(n)}$ denken wir uns um irgendeinen Punkt O des Systems als Zentrum zwei n -dimensionale Kugeln konstruiert; der Radius der einen sei die kleinste der Höhen des Elementarparallelepipedum, der Radius der anderen die kleinste der Längen seiner Kanten. Nach (3) kann in das Innere der ersten Kugel außer O kein weiterer Punkt des Systems fallen; dagegen liegen gewiß zwei solcher Punkte an den Enden eines bestimmten Durchmessers der zweiten, mithin jedenfalls nicht kleineren Kugel. Nach 1. oder 3. können wir alle Punkte bestimmen, welche in der Schicht zwischen den beiden Kugeln, die Begrenzungen mit eingerechnet, sich vorfinden; ihre Anzahl ist nach den dortigen Sätzen eine beschränkte. Unter diesen Punkten werden dann ein oder vielleicht mehrere Paare vorhanden sein, welche dem Punkte O am nächsten liegen. Die Entfernung dieser nächstgelegenen Punkte von O bezeichnen wir mit $\sqrt{M}$; wegen der Regelmäßigkeit des Punktsystems ist dieses dann überhaupt die kleinste Entfernung zweier Punkte, welche im Systeme vorkommt. Zugleich ist M die kleinste, von Null verschiedene Größe, welche durch die, zur gegebenen Anordnung des Systems gehörige quadratische Form f mittels ganzer Zahlen darstellbar ist; wir nennen M das *Minimum* dieser Form f . Fast evident erscheint nun die folgende wichtige Eigenschaft:

Die kleinste Entfernung zweier Punkte in einem regelmäßigen Punktsysteme kann nicht einen gewissen, durch die mittlere Dichtigkeit des Systems bestimmten Betrag übersteigen.

Denn denken wir uns um jeden Punkt des Systems einen n -dimensionalen Würfel von der Kante $\frac{1}{\sqrt{n}} \sqrt{M}$ abgegrenzt, indem wir jedesmal den Punkt als Mittelpunkt des Würfels nehmen — wir können uns etwa alle diese Würfel parallel orientiert vorstellen —, so sind die vom Mittelpunkte am weitesten abliegenden Punkte eines solchen Würfels jedesmal seine Eckpunkte, und die Entfernung dieser vom Mittelpunkte beträgt das $\frac{1}{2} \sqrt{n}$ -fache der Kante, also hier $\frac{1}{2} \sqrt{M}$. Wegen der Bedeutung der Länge $\sqrt{M}$ können daher diese Würfel sich niemals durchdringen, sie können höchstens unter Umständen in ihren Eckpunkten zusammentreffen, müssen im übrigen aber außerhalb ihrer Seitenflächen noch einen freien

Raum zwischen sich lassen. Ziehen wir diesen freien Raum in Betracht, so kommt also in einem, überallhin gleichmäßig ins Unendliche ausgedehnten Raume auf einen Raumteil gleich dem Volumen eines der Würfel im Durchschnitt weniger als ein Punkt des Systems. Dieses Volumen muß also nach den Betrachtungen in 2. kleiner sein als das Volumen des Elementarparallelepipedum, d. h. wir haben:

$$\left(\frac{1}{\sqrt[n]{n}} \sqrt[n]{M}\right)^n < \sqrt[n]{\Delta},$$

oder:

$$(6a) \quad M < n \sqrt[n]{\Delta},$$

womit unsere Behauptung erwiesen ist. In dem sehr einfachen Falle $n = 1$, den wir hier stillschweigend übergangen haben, müßte in diesen Formeln offenbar das Gleichheitszeichen statt $<$ genommen werden.

Wir haben so den folgenreichen Satz von Herrn Hermite über das Minimum einer positiven quadratischen Form in das rechte Licht gesetzt und zugleich in $n \sqrt[n]{\Delta}$ eine sehr viel engere Grenze für dieses Minimum gefunden, als sie, die kleinsten Zahlen n ausgenommen, bisher bekannt ist (s. unten 10.). Wir können aber sofort auch diese Grenze noch einschränken. Konstruieren wir nämlich um jeden Punkt des Systems als Mittelpunkt eine n -dimensionale Kugel mit dem Radius $\frac{1}{2} \sqrt[n]{M}$, so müssen auch diese, den vorhin konstruierten Würfeln umschriebenen Kugeln sich gegenseitig vollständig ausschließen und zwischen sich noch einen freien Raum lassen, und es muß also auch das Volumen einer solchen Kugel kleiner sein als das Volumen des Elementarparallelepipedum. Nun beträgt das Volumen einer n -dimensionalen Kugel vom Radius 1 bekanntlich:

$$\frac{\left\{\Gamma\left(\frac{1}{2}\right)\right\}^n}{\Gamma\left(1 + \frac{n}{2}\right)},$$

d. h. je nachdem n gerade oder ungerade ist:

$$\frac{\pi^{\frac{n}{2}}}{1 \cdot 2 \cdot 3 \cdots \frac{n}{2}} \quad \text{oder} \quad \frac{\frac{n+1}{2} \cdot \pi^{\frac{n-1}{2}}}{1 \cdot 3 \cdot 5 \cdots n};$$

also finden wir:

$$(6b) \quad \frac{\left\{\Gamma\left(\frac{1}{2}\right)\right\}^n}{\Gamma\left(1 + \frac{n}{2}\right)} \left(\frac{1}{2} \sqrt[n]{M}\right)^n < \sqrt[n]{\Delta}.$$

Durch Benutzung des asymptotischen Ausdrucks der Γ -Funktion folgt daraus leicht:

$$M < \frac{2n}{\pi e} \sqrt[n]{\frac{1}{n \pi e^{\frac{1}{3n}} \sqrt[n]{\Delta}}},$$

so daß diese zweite Grenze für das Minimum bei großen Werten von n ungefähr das $\frac{2}{\pi e} = 0,234 \dots$ -fache der früher gefundenen ausmacht. Später werden wir noch engere Grenzen für das Minimum kennen lernen.

7. Anwendung auf die Theorie der algebraischen Zahlen.

Eine der ersten Anwendungen, welche Herr Hermite von der Existenz einer Grenze für das Minimum positiver quadratischer Formen gemacht hat, betraf die Theorie der algebraischen Zahlen. Die großen Fortschritte, welche auf diesem Gebiete seitdem erzielt sind, und andererseits die im vorhergehenden gefundene natürlichere Grenze für das Minimum ermöglichen es uns, diese Anwendung wesentlich zu vertiefen und sie zugleich in vollkommenerer Form zur Darstellung zu bringen.

Es sei θ eine Wurzel einer irreduktiblen ganzzahligen Gleichung von einem Grade n , welcher größer als Eins sei; und es bedeute $\mathfrak{o}$ das System aller ganzen algebraischen Zahlen, welche unter den rationalen Funktionen von θ mit ganzzahligen Koeffizienten überhaupt zu finden sind. Es sei ferner $\omega_1, \omega_2, \dots, \omega_n$ irgendeine Reihe von n Zahlen aus $\mathfrak{o}$, für welche das Quadrat der Determinante

$$|\omega_k^{(h)}| \quad (h, k = 1, 2, \dots, n)$$

aus den n konjugierten Reihen von Null verschieden und dazu dem absoluten Betrage nach möglichst klein ausfalle, und der dabei eintretende Wert dieses Quadrats, die sogenannte *Diskriminante* von $\mathfrak{o}$, heiße D . Das System $\mathfrak{o}$ stimmt dann genau überein mit den Werten der Form

$$\omega_1 x_1 + \omega_2 x_2 + \dots + \omega_n x_n = \omega$$

für alle möglichen ganzen Zahlen $x_1, x_2, \dots, x_n$, und diese Werte sind untereinander alle verschieden*). Wenden wir dieselben Zeichen $x_1, x_2, \dots, x_n$ für die Unbestimmten einer beliebigen, wesentlich positiven Form f mit entsprechender Zahl n an, so kann daher ein, dieser Form f gemäß parallelepipedisch aufgebautes regelmäßiges Punktsystem $\mathfrak{D}$ gewissermaßen

*) Kronecker, Grundzüge einer arithmetischen Theorie der algebraischen Größen. Festschrift zu Herrn Kummers Doktorjubiläum. Crelles Journal Bd. 92, S. 99. (Werke, Bd. II, S. 360.) — Dedekind, Allgemeine Zahlentheorie (Suppl. XI zu den Vorlesungen über Zahlentheorie von Dirichlet, III. [[oder IV.]] Aufl.). — Für unsern speziellen Zweck liegen die Dedekindschen Begriffsbestimmungen besonders günstig; auf das soeben genannte Werk beziehen sich im folgenden die Zitate mit dem Buchstaben D .

als Träger des gesamten Zahlensystems $\mathfrak{o}$ betrachtet werden; wir haben nur festzusetzen, welcher Punkt der Zahl $\omega = 0$ entsprechen soll.

Unter einem *Ideal* des Gebietes $\mathfrak{o}$ versteht man nach Herrn Dedekind jedes in $\mathfrak{o}$ enthaltene und nicht aus der Zahl Null allein bestehende Zahlensystem $\mathfrak{a}$, dessen Inhalt keine Bereicherung erfahren könnte, weder wenn man Summen und Differenzen aus seinen Zahlen, noch wenn man Produkte aus seinen Zahlen in Zahlen aus $\mathfrak{o}$ hinzunehmen wollte (D. § 168 [[IV. Aufl., § 177]]). Als Träger eines Ideals $\mathfrak{a}$ erscheint ein, im Punktsysteme $\mathfrak{D}$ im Sinne von 5. enthaltenes, ebenfalls parallelepipedischer Anordnungen fähiges, regelmäßiges n -dimensionales Punktsystem $\mathfrak{A}$; der Quotient aus der mittleren Dichtigkeit des Punktsystems $\mathfrak{D}$ und der mittleren Dichtigkeit dieses darin enthaltenen Punktsystems $\mathfrak{A}$ heißt die *Norm* des Ideals $\mathfrak{a}$, in Zeichen: $Nm(\mathfrak{a})$. Es gibt immer nur eine beschränkte Anzahl von Idealen, welche dieselbe Norm haben. Das System $\mathfrak{o}$, selbst ein Ideal, ist offenbar das einzige von der Norm 1. Die Gesamtheit aller Zahlen in $\mathfrak{o}$, welche durch eine bestimmte, von Null verschiedene Zahl η aus $\mathfrak{o}$ teilbar sind, konstituiert ein sogenanntes *Hauptideal* $\mathfrak{o}\eta$; die Norm eines solchen ist der absolute Wert der Norm von η , d. i. des Produkts der n konjugierten Zahlen $\eta', \eta'', \dots, \eta^{(n)}$, welche zu den einzelnen n Wurzeln der irreduktiblen Ausgangsgleichung in derselben Beziehung stehen wie die Zahl η zu der Wurzel θ dieser Gleichung.

Unter dem *Produkte* $\mathfrak{a}\mathfrak{b}$ zweier Ideale $\mathfrak{a}$ und $\mathfrak{b}$ versteht man den Inbegriff aller Zahlen, welche sich als ein Produkt aus einer Zahl in $\mathfrak{a}$ und einer Zahl in $\mathfrak{b}$ oder als Summe mehrerer solcher Produkte darstellen lassen; das Produkt $\mathfrak{a}\mathfrak{b}$ ist wieder ein Ideal und seine Norm das Produkt der Normen von $\mathfrak{a}$ und von $\mathfrak{b}$ (D. § 170 [[IV. Aufl., § 177 u. § 180]]). Beziehungen zwischen Produkten aus Idealen lassen ganz analoge Folgerungen zu wie Beziehungen zwischen Produkten aus rationalen ganzen Zahlen; das Ideal $\mathfrak{o}$ spielt dabei die Rolle der Zahl 1.

Zu jeder von Null verschiedenen Zahl μ eines Ideals $\mathfrak{a}$ gibt es ein bestimmtes Ideal $\mathfrak{m}$, welches die Gleichung $\mathfrak{o}\mu = \mathfrak{a}\mathfrak{m}$ befriedigt und also die Fähigkeit besitzt, durch sein Hinzutreten als Faktor das Ideal $\mathfrak{a}$ in ein Hauptideal zu verwandeln (D. § 175 [[IV. Aufl., § 178]]). Die Ideale werden nach den Multiplikatoren klassifiziert, welche geeignet sind, sie in Hauptideale zu verwandeln und über diese Multiplikatoren wollen wir nun einen wichtigen Satz ableiten. Zu dem Ende legen wir jedoch eine quadratische Form f von besonderer Beschaffenheit zugrunde, nämlich wir setzen:

$$f = \sum_h \lambda_h (\text{abs. } \omega_1^{(h)} x_1 + \omega_2^{(h)} x_2 + \dots + \omega_n^{(h)} x_n)^2 \quad (h = 1, 2, \dots, n);$$

die linearen Formen in diesem Ausdrucke sollen die n mit der Form ω

konjugierten Formen vorstellen; unter $(\text{abs.})^2$ soll das Quadrat des absoluten Betrags einer solchen Form verstanden werden, die Variablen als reelle Größen gedacht; ferner sollen die λ_h beliebige positive Konstanten bedeuten. Ein solches f ist eine wesentlich positive quadratische Form, und die Determinante dieser Form hat den Ausdruck $\prod_h \lambda_h \text{ abs. } D$, unter $\text{abs. } D$ den absoluten Wert der Diskriminante D verstanden. Die mittlere Dichtigkeit in dem, zu einem Ideal $\mathfrak{a}$ gehörigen Punktsysteme $\mathfrak{A}$ wird demnach

$$1 : \text{Nm}(\mathfrak{a}) \sqrt{\prod_h \lambda_h \cdot \text{abs. } D}$$

betragen. Fassen wir nun in dem Punktsysteme $\mathfrak{A}$ einen Punkt ins Auge, welcher möglichst nahe dem Nullpunkte liegt, und benutzen wir die in (6a) gegebene Grenze für die kleinste Entfernung zweier Punkte in einem regelmäßigen Punktsysteme, so können wir aus dem Orte dieses Punktes n Zahlen $x_1, x_2, \dots, x_n$ erschließen, für welche

$$\omega_1 x_1 + \omega_2 x_2 + \dots + \omega_n x_n = \mu$$

eine Zahl in $\mathfrak{a}$ ist, und zugleich erweist sich für diese Zahlen der Ausdruck

$$\sum_h \lambda_h (\text{abs. } \mu^{(h)})^2 < n \sqrt{\prod_h \lambda_h \cdot (\text{Nm } \mathfrak{a})^2 \text{ abs. } D}. \quad (h = 1, 2, \dots, n)$$

Ein besonderer Nachdruck ist aus einem bald ersichtlichen Grunde darauf zu legen, daß hier das Zeichen $<$ und nicht etwa $\leq$ sich einfindet. Benutzen wir nun, daß eine Summe von n positiven Größen niemals kleiner ist als das n -fache der n^{ten} Wurzel aus dem Produkte der n Größen, und setzen wir zugleich $(\text{Nm } \mu)^2$ für $\prod_h (\text{abs. } \mu^{(h)})^2$, so können wir aus der vorstehenden Ungleichung die weitere entnehmen:

$$n \sqrt{\prod_h \lambda_h \cdot (\text{Nm } \mu)^2} < n \sqrt{\prod_h \lambda_h \cdot (\text{Nm } \mathfrak{a})^2 \text{ abs. } D}. \quad (h = 1, 2, \dots, n)$$

Ist $\mathfrak{m}$ das Ideal, welches die Gleichung $\mathfrak{o}\mu = \mathfrak{a}\mathfrak{m}$ befriedigt, so haben wir $\text{Nm}(\mathfrak{a})\text{Nm}(\mathfrak{m}) = \pm \text{Nm}(\mu)$, und wir finden demnach:

$$\text{Nm}(\mathfrak{m}) < \sqrt{\text{abs. } D}.$$

Zu jedem Ideal gibt es behufs Herstellung eines Hauptideals mindestens einen Multiplikator, bei welchem die Norm weniger beträgt als die Wurzel aus dem absoluten Werte der Diskriminante.

Als eine spezielle Folgerung geht daraus der bekannte Satz hervor, daß eine endliche Anzahl von Multiplikatoren ausreichend ist, um alle Ideale in Hauptideale zu verwandeln*). Eine andere, sehr bemerkenswerte

*) Dieser Satz ist auf Herrn Kronecker zurückzuführen. Vgl. die Bemerkungen auf S. 64 der Festschrift zu Herrn Kummers Doktorjubiläum, Crelles Journal, Bd. 92. (Werke, Bd. II, S. 320.)

Folgerung ist diese: Da die Norm eines Ideals eine ganze Zahl, mindestens gleich Eins ist, so ergibt die letzte Ungleichung $1 < \sqrt{\text{abs. } D}$, also muß D von ± 1 verschieden sein, d. h.:

Jede Diskriminante enthält Primzahlen als Faktoren.

Diese das Wesen der algebraischen Zahlen tief berührende Eigenschaft findet sich auf Seite 21 der eben zitierten Festschrift von Herrn Kronecker ausgesprochen; doch ist ein Beweis dieser Eigenschaft bisher nicht veröffentlicht worden.

Überhaupt kommen die kleinsten positiven wie negativen Zahlen bis zu gewissen von der jedesmaligen Ordnung n abhängigen Grenzen als Werte von Diskriminanten nicht vor.

Denn benutzen wir die zweite in 6. gefundene Grenze für das Minimum einer positiven quadratischen Form, so finden wir im übrigen nach genau demselben Verfahren, wie bei der ersten Grenze, eine schärfere Ungleichung, nämlich die folgende:

$$(7b) \quad \text{Nm}(m) < \frac{\Gamma\left(1 + \frac{n}{2}\right)}{\left\{\Gamma\left(\frac{1}{2}\right)\right\}^n} \frac{2^n}{n^{\frac{n}{2}}} \sqrt{\text{abs. } D},$$

und diese können wir wieder mit der Ungleichung $1 \leq \text{Nm}(m)$ verbinden; so zeigt sich, daß der absolute Wert einer Diskriminante n^{ter} Ordnung

sicherlich immer die GröÙe $\frac{\left(\frac{\pi e}{2}\right)^n}{n \pi e^{\frac{1}{3n}}}$ übertrifft.

Die zuletzt gefundene Grenze für die Norm eines Multiplikators wollen wir noch in einem Beispiele anwenden. Herr Wolfskehl*) hat vor kurzem den Nachweis geliefert, daß der zweite Faktor der Klassenanzahl für die aus den 13^{ten} Wurzeln der Einheit gebildeten Zahlen gleich Eins ist. Dieser Nachweis erforderte außer einer Benutzung der Reuschleichen Tafeln noch recht verwickelte Rechnungen. Wir können hier einfacher zu demselben Satze gelangen und noch hinzufügen, daß die gleiche Erscheinung auch bei den 17^{ten} und 19^{ten} Wurzeln der Einheit eintritt; ja später werden wir sogar durch ähnliche Mittel dieses Resultat noch weiter auszudehnen imstande sein.

Stellt λ eine ungerade Primzahl vor, so bedeutet der zweite Faktor der Klassenanzahl für die aus den λ^{ten} Einheitswurzeln gebildeten Zahlen — wir machen augenblicklich von der Kummerschen Terminologie Gebrauch — dasselbe wie die Klassenanzahl für die aus den $\frac{\lambda-1}{2}$ zwei-

*) Crelles Journal, Bd. 99, S. 173.

gliedrigen Perioden dieser Einheitswurzeln gebildeten Zahlen. Die Diskriminante des Systems dieser letzteren Zahlen hat den Ausdruck $\lambda^{\frac{1}{2}(\lambda-3)}$; setzen wir in die Ungleichung (7b) diese Größe für D und zugleich $\frac{\lambda-1}{2}$ für n , so erlangt die rechte Seite dort folgende Werte:

$$\text{für } \lambda = 5, \quad 7, \quad 11, \quad 13, \quad 17, \quad 19, \\ 1, \dots, 2, \dots, 13, \dots, 34, \dots, 311, \dots, 1027, \dots$$

Bis zu diesen Grenzen hätten wir also höchstens die Normen der Multiplikatoren zur Hervorbringung wirklicher Zahlen zu suchen, und wenn alle Zahlen bis zu diesen Grenzen sich in wirkliche Faktoren zerlegen lassen sollten, so kommen in den hier betrachteten Fällen ideale Multiplikatoren und demgemäß auch ideale Zahlen überhaupt nicht vor. Nun entnehmen wir aus den Reuschleschen Tafeln, daß für die soeben aufgezählten Werte von λ alle Zahlen unter 1000 sich in wirkliche Faktoren zerlegen lassen. Wir haben also nur noch in bezug auf $\lambda = 19$ festzustellen, daß hier die Primzahlen zwischen 1000 und 1027, das sind 1009, 1013, 1019, 1021, ebenfalls einer solchen Zerlegung innerhalb des durch die entsprechenden zweigliedrigen Perioden bestimmten Gebiets fähig sind. Nun gehören in bezug auf die Primzahl 19 die Zahlen 1009 und 1021 zum Exponenten 18, die Zahl 1013 zum Exponenten 9; diese Zahlen sind also in dem fraglichen Zahlengebiet der zweigliedrigen Perioden selbst noch Primzahlen. Die Zahl 1019 endlich gehört modulo 19 zum Exponenten 6, ihre Primfaktoren werden also von den drei sechsgliedrigen Perioden der 19^{ten} Einheitswurzeln abhängen. Diese sind die Wurzeln der Gleichung $\eta^3 + \eta^2 - 6\eta - 7 = 0$, deren Diskriminante den Wert $D = 19^2$ hat. Da nun in dem durch diese Wurzeln bestimmten Gebiete nach den Reuschleschen Tafeln die Zahlen bis zu $\sqrt{D} = 19$ in wirkliche Faktoren zerlegbar sind, so können in diesem Gebiete ideale Teiler nicht existieren, und demnach muß auch 1019 in drei wirkliche Faktoren zerlegt werden können.

IX.

Théorèmes arithmétiques.

Extrait d'une lettre de M. H. Minkowski à M. Hermite.

(Comptes rendus de l'Académie des Sciences, t. 112, pp. 209—212.)

» La méthode géométrique de mon travail*), traduite en langue purement analytique, conduit à ce théorème susceptible d'une application très étendue:

» Soit n un nombre plus grand que 1; soient $\xi, \eta, \zeta, \dots, n$ formes linéaires indépendantes à n variables $x, y, z, \dots$. Parmi ces formes, soient β paires d'imaginaires conjuguées et les autres $n - 2\beta = \alpha$ formes réelles. L'un ou l'autre des nombres α et β peut aussi être égal à zéro. Soit Δ le déterminant des formes $\xi, \eta, \zeta, \dots$. Soit enfin p une quantité quelconque ≥ 1 . On peut toujours assigner à $x, y, z, \dots$ des valeurs entières, de sorte que la somme

$$(\text{abs. } \xi)^p + (\text{abs. } \eta)^p + (\text{abs. } \zeta)^p + \dots$$

soit différente de zéro et en même temps plus petite que la quantité

$$\left\{ \left(\frac{2}{\pi} \right)^\beta \frac{\Gamma \left(1 + \frac{n}{p} \right)}{\left[\Gamma \left(1 + \frac{1}{p} \right) \right]^\alpha 2^{-\frac{2\beta}{p}} \left[\Gamma \left(1 + \frac{2}{p} \right) \right]^\beta} \text{abs. } \Delta \right\}^{\frac{p}{n}},$$

qui est elle-même plus petite que

$$n (\text{abs. } \Delta)^{\frac{p}{n}}.$$

Ici abs. signifie « valeur absolue de » et Γ désigne la fonction gamma.

» En suivant une voie indiquée dans vos admirables lettres à Jacobi, je tirerai du théorème que je viens d'exposer plusieurs conclusions fondamentales sur les nombres algébriques.

*) Über die positiven quadratischen Formen und über kettenbruchähnliche Algorithmen (Journal de Crelle, t. 107, p. 278). Diese Ges. Abhandlungen, Bd. I, S. 243—260.

» Soit un corps algébrique quelconque, irréductible et d'ordre n , et soit ξ une forme linéaire qui, pour toutes les valeurs entières de ses n variables $x, y, z, \dots$, représente tous les entiers algébriques de ce corps; soient, de plus, $\eta, \zeta, \dots$ les $n - 1$ formes conjuguées à ξ . Le discriminant du corps est représenté par le carré du déterminant Δ , et ce carré est un entier rationnel D du signe $(-1)^\beta$. En faisant usage de l'inégalité

$$(\text{abs. } \xi \eta \zeta \dots)^p \leq \left[\frac{(\text{abs. } \xi)^p + (\text{abs. } \eta)^p + (\text{abs. } \zeta)^p + \dots}{n} \right]^n,$$

et en remarquant que $\text{abs. } \xi \eta \zeta \dots$ est un entier ≥ 1 , pourvu que $x, y, z, \dots$ soient des entiers et qu'ils ne s'évanouissent pas tous, les inégalités du théorème énoncé entraîneront celles-ci:

$$1 < \left\{ \left(\frac{2}{\pi} \right)^\beta \frac{n^{-\frac{n}{p}} \Gamma \left(1 + \frac{n}{p} \right)}{\left[\Gamma \left(1 + \frac{1}{p} \right) \right]^\alpha 2^{-\frac{2\beta}{p}} \left[\Gamma \left(1 + \frac{2}{p} \right) \right]^\beta} \right\}^2 \text{abs. } D < \text{abs. } D.$$

» Faisant d'abord abstraction du terme intermédiaire, nous avons ainsi démontré le postulat profond de M. Kronecker*), que chaque discriminant est différent de ± 1 , c'est-à-dire que *chaque discriminant contient des nombres premiers comme facteurs*. C'est là un détail bien digne d'attention. Tout nombre algébrique irrationnel a ainsi ses nombres premiers critiques, comme toute fonction algébrique irrationnelle a ses points d'embranchement.

» Le terme dont nous n'avons pas tenu compte nous fournit pour la valeur absolue d'un discriminant des limites inférieures plus complètes. Ces autres limites, où figure encore le nombre β , s'accroissant indéfiniment avec l'ordre n , il est évident qu'un nombre donné quelconque ne peut être discriminant que pour un nombre fini d'ordres n .

» De quelle manière fixera-t-on le mieux la quantité p , assujettie jusqu'à présent à la seule condition de ne pas être moindre que l'unité? On se convaincra aisément que les limites dont nous venons de parler devront s'agrandir aussi longtemps que la valeur de p décroît. Ce n'est donc pas quand p est égal à 2, valeur qui répond aux formes quadratiques, mais dans le cas de $p = 1$, que ces limites seront le plus avancées. Il en résulte enfin ce théorème:

» *Le discriminant d'un corps algébrique, faisant partie de n corps conjugués dont 2β sont imaginaires et $n - 2\beta$ réels, est en valeur absolue toujours plus grand que*

$$\left[\left(\frac{\pi}{4} \right)^\beta \frac{n^n}{2 \cdot 3 \dots n} \right]^2.$$

*) Journal de Crelle, t. 92, (Werke, Bd. II, S. 269).

» Par exemple, un discriminant du deuxième ordre doit être ou > 4 ou < -2 , Les valeurs les plus petites 5 et -3 se trouvent dans les équations $\omega^2 + \omega - 1 = 0$ et $\omega^2 + \omega + 1 = 0$.

» Un discriminant du troisième ordre doit être ou > 20 , ... ou < -12 , De la limite précise du minimum des formes quadratiques positives ternaires on aurait tiré, en suivant une marche tout analogue, les inégalités $D \geq 13,5$ ou $\leq -13,5$. La limite que nous avons trouvée plus haut n'est donc pas, il est vrai, une limite précise, mais malgré cela elle nous fournit déjà des résultats que les formes quadratiques n'ont pas encore donnés.»

X.

Über Geometrie der Zahlen.

Bericht über einen Vortrag zu Halle.

(Verhandlungen der 64. Naturforscher- und Ärzteversammlung zu Halle, 1891, S. 13
und

Jahresbericht der Deutschen Mathematiker-Vereinigung, Band 1, S. 64—65.)

Wenn man für den Raum rechtwinklige Koordinaten einführt, so entsprechen den Systemen von drei *ganzen* Zahlen diskrete Punkte, welche derart über den Raum verstreut liegen, daß sie eine gewisse Nähe in bezug auf jede beliebige Raumstelle erreichen. Den Inbegriff aller dieser Punkte mit lauter Koordinaten, die ganze Zahlen sind, nennt der Vortragende das dreidimensionale *Zahlengitter*; unter dem Titel „*Geometrie der Zahlen*“ begreift er geometrische Studien über das dreidimensionale Zahlengitter und über das entsprechende Gebilde in der Ebene, und in weiterem Sinne auch die Ausdehnung der Ergebnisse solcher Studien auf Mannigfaltigkeiten beliebiger Ordnung. Natürlich besitzt jede Aussage über die Zahlengitter einen rein arithmetischen Kern. Das Wort „Geometrie“ erscheint aber durchaus am Platze im Hinblick auf Fragestellungen, zu welchen die geometrische Anschauung verhilft, und auf Untersuchungsmethoden, welche fortwährend durch geometrische Begriffe ihre Richtung angewiesen erhalten.

Der Vortragende hat sich in betreff der Zahlengitter hauptsächlich zwei Fragen gestellt; sie ergänzen einander in gewisser Beziehung, und folgendes ist ihnen gemeinsam: Es handelt sich, wenn speziell vom Raume gesprochen wird, jedesmal um eine sehr allgemeine Kategorie von Körpern, welche so konstruiert werden, daß sie einen bestimmten Punkt des Zahlengitters — es sei dies etwa der Nullpunkt — in gewisser Weise umschließen, und es soll dann jedesmal bei diesen Körpern eine gewisse Eigenschaft in bezug auf das Zahlengitter allein durch die Größe des Inhalts der Körper zustande kommen.

Die erste Kategorie von Körpern besteht aus allen denjenigen Körpern, welche im Nullpunkte einen Mittelpunkt haben, und deren Begrenzung nach außen hin nirgends konkav ist; und die fragliche Eigen-

schaft für diese Kategorie lautet: Wenn der Inhalt eines Körpers dieser Kategorie $\geq 2^3$ ist, so schließt der Körper notwendig noch weitere Punkte des Zahlengitters außer dem Nullpunkte ein.

Die zweite Kategorie von Körpern ist noch umfassender; sie besteht aus allen Körpern, welche den Nullpunkt enthalten, und deren Oberfläche, vom Nullpunkte aus gesehen, nach jeder Richtung hin nur einen Punkt darbietet; und die fragliche Eigenschaft für diese zweite Kategorie lautet: Wenn der Inhalt eines Körpers dieser Kategorie

$$\leq 1 + \frac{1}{2^3} + \frac{1}{3^3} + \frac{1}{4^3} + \dots$$

ist, so können stets Deformationen des Körpers angegeben werden, bei welchen der Inhalt sich nicht ändert, der Nullpunkt fest bleibt und gerade Linien gerade Linien bleiben, und nach deren Ausführung alle Punkte des Zahlengitters mit Ausnahme des Nullpunkts ihren Ort außerhalb des Körpers finden.

Der Vortragende weist auf die außerordentliche Tragweite dieser, in ihrer Allgemeinheit ebenso einfach wie plausibel klingenden Sätze hin.

XI.

Extrait d'une lettre adressée à M. Hermite.

(Bulletin des Sciences mathématiques, 2^e série, t. XVII, pp. 24—29.)

Veillez bien permettre, Monsieur, que je vous donne un rapide résumé de mon Ouvrage.*) La plus grande partie du livre traite des fonctions φ à n variables $x_1, x_2, \dots, x_n$, qui, comme la racine carrée d'une forme quadratique positive, satisfont aux conditions

$$\begin{aligned} (A) \quad & \left\{ \begin{array}{l} \varphi(x_1, x_2, \dots, x_n) > 0, \quad \text{si l'on n'a pas } x_1 = 0, \quad x_2 = 0, \dots, x_n = 0, \\ \varphi(0, 0, \dots, 0) = 0, \\ \varphi(tx_1, tx_2, \dots, tx_n) = t\varphi(x_1, x_2, \dots, x_n), \quad \text{si } t > 0, \end{array} \right. \\ (B) \quad & \varphi(x_1 + y_1, x_2 + y_2, \dots, x_n + y_n) \leq \varphi(x_1, x_2, \dots, x_n) + \varphi(y_1, y_2, \dots, y_n), \\ (C) \quad & \varphi(-x_1, -x_2, \dots, -x_n) = \varphi(x_1, x_2, \dots, x_n). \end{aligned}$$

Soient $\xi_1, \xi_2, \dots, \xi_v$ un nombre fini de formes linéaires à coefficients réels et aux variables $x_1, x_2, \dots, x_n$, et parmi ces formes soient n formes à déterminant différent de zéro. Soit

$$\Phi(x_1, x_2, \dots, x_n)$$

le maximum parmi les valeurs absolues de $\xi_1, \xi_2, \dots, \xi_v$. Une telle fonction Φ satisfera évidemment aux conditions d'une fonction φ .

J'établis d'abord ce théorème:

φ étant une solution quelconque de (A), (B), (C), et δ une quantité positive choisie à volonté, on peut toujours trouver des fonctions Φ comme je viens de les caractériser, de sorte que, pour toutes les valeurs possibles de $x_1, x_2, \dots, x_n$, on ait

$$1 \leq \frac{\varphi}{\Phi} < 1 + \delta.$$

Il en résulte que l'intégrale $\int \int \dots \int dx_1 dx_2 \dots dx_n$ étendue sur le domaine $\varphi(x_1, x_2, \dots, x_n) \leq 1$ aura toujours une valeur déterminée. Soit J cette valeur. Je démontre alors que l'on peut toujours trouver des nombres entiers $x_1, x_2, \dots, x_n$ pour lesquels on ait

$$(I) \quad 0 < \varphi(x_1, x_2, \dots, x_n) \leq \frac{2}{\sqrt[n]{J}}.$$

*) Gemeint ist die „Geometrie der Zahlen.“ (Anm. d. Herausg.)

J'ajoute ce théorème supplémentaire :

Le cas qu'il n'existe pas de nombres entiers $x_1, x_2, \dots, x_n$, pour lesquels on ait

$$0 < \varphi(x_1, x_2, \dots, x_n) < \frac{2}{\sqrt[n]{J}},$$

ne peut se présenter que chez des fonctions Φ provenant d'un nombre ν de formes linéaires $\leq 2^n - 1$.

La plus simple application du théorème (I) est la suivante :

$\xi_1, \xi_2, \dots, \xi_n$ étant n formes linéaires à coefficients réels quelconques et à déterminant égal à ± 1 , on peut toujours donner à $x_1, x_2, \dots, x_n$ des valeurs entières qui ne s'évanouissent pas toutes et de sorte que les valeurs absolues de $\xi_1, \xi_2, \dots, \xi_n$ soient toutes ≤ 1 .

L'énoncé plus exact de ce théorème est que l'on peut trouver des nombres entiers $x_1, x_2, \dots, x_n$, qui ne s'évanouissent pas tous, et de sorte que les valeurs absolues de $\xi_1, \xi_2, \dots, \xi_n$ soient toutes < 1 , excepté le cas où les formes $\xi_1, \xi_2, \dots, \xi_n$, par une substitution linéaire à coefficients entiers et à déterminant ± 1 , peuvent être transformées de manière que, abstraction faite de l'ordre, elles deviennent

$$x_1, \quad a_{21}x_1 + x_2, \quad \dots, \quad a_{n1}x_1 + a_{n2}x_2 + \dots + x_n.$$

Ainsi, par exemple, $a_1, a_2, \dots, a_{n-1}$ étant des quantités réelles quelconques et T une quantité > 1 , il y aura des nombres entiers $x_1, x_2, \dots, x_{n-1}, x_n$, parmi lesquels x_n est différent de zéro, de sorte que les valeurs absolues de

$$x_1 - a_1 x_n, \quad x_2 - a_2 x_n, \quad \dots, \quad x_{n-1} - a_{n-1} x_n, \quad \frac{x_n}{T^n}$$

soient toutes $< \frac{1}{T}$, excepté le cas où T est un nombre entier, et $a_1, a_2, \dots, a_{n-1}$, abstraction faite de l'ordre, ont des expressions $\frac{T_1}{T}, \frac{T_2}{T^2}, \dots, \frac{T_{n-1}}{T^{n-1}}$, dans lesquelles $T_1, T_2, \dots, T_{n-1}$ sont des nombres entiers premiers à T .

En appliquant le théorème (I) à la fonction Φ qui est définie à l'aide des $2n - 2$ formes

$$x_1 - a_1 x_n \pm \frac{x_n}{S}, \quad x_2 - a_2 x_n \pm \frac{x_n}{S}, \quad \dots, \quad x_{n-1} - a_{n-1} x_n \pm \frac{x_n}{S},$$

on conclut que l'on peut toujours trouver des nombres entiers $x_1, x_2, \dots, x_{n-1}, x_n$, sans diviseur commun et parmi lesquels x_n est positif, de sorte que les valeurs absolues de

$$\frac{x_1}{x_n} - a_1, \quad \frac{x_2}{x_n} - a_2, \quad \dots, \quad \frac{x_{n-1}}{x_n} - a_{n-1}$$

soient plus petites qu'une quantité positive ε choisie à volonté, et en même temps

$$< \frac{n-1}{n \frac{n}{n-1}}.$$

A l'aide de certaines autres fonctions φ , on obtient les théorèmes analogues pour les quantités complexes.

Les théorèmes que j'ai exposés dans ma dernière lettre sont aussi des conséquences spéciales du théorème (I). De ce théorème découlent enfin les théorèmes de Dirichlet sur les unités complexes.

Ensuite j'établis cette généralisation du théorème (I):

Pour toute fonction φ , satisfaisant aux conditions (A), (B), (C), on peut trouver n^2 nombres entiers l_{hk} à déterminant différent de zéro, de sorte que l'on ait

$$\varphi(l_{11}, l_{21}, \dots, l_{n1}) \varphi(l_{12}, l_{22}, \dots, l_{n2}) \dots \varphi(l_{1n}, l_{2n}, \dots, l_{nn}) \leq \frac{2^n}{J}.$$

Le déterminant l_{hk} sera alors toujours $\leq 1 \cdot 2 \dots n$.

Je fais aussi quelques remarques sur les cas extrêmes de cette relation. Des applications de ce théorème, je ne cite ici que ces deux:

Soient $a_{hk}(h, k = 1, 2, \dots, n)$ n^2 quantités réelles à déterminant différent de zéro, et soit D la valeur absolue de ce déterminant. Il y aura ou n^2 nombres entiers l_{hk} à déterminant différent de zéro, de sorte que le système composé

$$\begin{pmatrix} a_{11} & \dots & a_{1n} \\ \dots & \dots & \dots \\ a_{n1} & \dots & a_{nn} \end{pmatrix} \begin{pmatrix} l_{11} & \dots & l_{1n} \\ \dots & \dots & \dots \\ l_{1n} & \dots & l_{nn} \end{pmatrix} = \begin{pmatrix} b_{11} & \dots & b_{1n} \\ \dots & \dots & \dots \\ b_{n1} & \dots & b_{nn} \end{pmatrix}$$

satisfasse à toutes les n^n inégalités

$$\pm b_{h_1 1} b_{h_2 2} \dots b_{h_n n} < D, \quad (h_1 = 1, 2, \dots, n; h_2 = 1, 2, \dots, n; \dots; h_n = 1, 2, \dots, n),$$

ou n^2 nombres entiers l_{hk} à déterminant ± 1 , de sorte que ce système composé, après une permutation convenable des lignes, prenne une forme

$$\begin{pmatrix} c_{11} & \dots & c_{1n} \\ \dots & \dots & \dots \\ c_{n1} & \dots & c_{nn} \end{pmatrix},$$

satisfaisant aux conditions

$$\begin{aligned} c_{hk} &= 0, & h > k, \\ 0 &< c_{11} \leq c_{22} \leq \dots \leq c_{nn}, \\ 0 &\leq c_{hk} < c_{hh}, & h < k. \end{aligned}$$

Une forme quadratique positive à n variables et à déterminant D peut toujours, par une substitution à coefficients entiers et à déterminant différent

de zéro $\left[\text{dont la valeur absolue est} < 2^n \frac{\Gamma\left(1 + \frac{n}{2}\right)}{\left(\Gamma\left(\frac{1}{2}\right)\right)^n} \right]$, être transformée en

une forme $\sum b_{hk} y_h y_k$, satisfaisant aux conditions

$$0 < b_{11} \leq b_{22} \leq \dots \leq b_{nn}, \quad \pm 2 b_{hk} \leq b_{hh}, \quad (h < k),$$

$$b_{11} b_{22} \dots b_{nn} < \left[\frac{2^n \Gamma\left(1 + \frac{n}{2}\right)}{\left(\Gamma\left(\frac{1}{2}\right)\right)^n} \right]^2 D.$$

Dans un autre Chapitre, je donne une nouvelle démonstration des théorèmes que Kronecker a établis dans son Mémoire *Näherungsweise ganzzahlige Auflösung linearer Gleichungen* (Werke, Bd. III, 1, S. 47).

Enfin je procède à la méthode de réduction des formes quadratiques positives que vous avez donnée en dernier lieu dans vos lettres à Jacobi. Il résulte de vos développements que l'on peut transformer toute forme quadratique positive f à n variables par une substitution à coefficients entiers et à déterminant ± 1 en une forme $\sum b_{hk} y_h y_k$, qui satisfait à toutes les inégalités

$$\sum b_{hk} p_h p_k \geq b_{mm},$$

où m est un des nombres $1, 2, \dots, n$ et où $p_1, p_2, \dots, p_n$ sont des nombres entiers quelconques, pour lesquels le plus grand diviseur commun de $p_m, p_{m+1}, \dots, p_n$ est égal à 1.

Je démontre que parmi ces inégalités on trouve un nombre fini d'où dérivent toutes les autres.

Pour une forme donnée f , il y aura, en général, 2^{n-1} formes satisfaisant à ces inégalités, et qui se déduiront d'une seule par les 2^n substitutions

$$y_1 = \pm z_1, \quad y_2 = \pm z_2, \quad \dots, \quad y_n = \pm z_n.$$

Ces inégalités étant linéaires dans les coefficients b_{hk} , on en conclut que, pour la somme

$$\chi(D+1) + \chi(D+2) + \dots + \chi(D+d),$$

$\chi(\Delta)$ désignant le nombre des classes de formes à coefficients entiers et à déterminant Δ , il existe une expression asymptotique

$$\gamma D^{\frac{n-1}{2}} d.$$

Je démontre que l'on a

$$\gamma = \frac{\Gamma\left(\frac{2}{2}\right) \Gamma\left(\frac{3}{2}\right) \dots \Gamma\left(\frac{n}{2}\right)}{\left[\Gamma\left(\frac{1}{2}\right)\right]^{2+3+\dots+n}} S_2 S_3 \dots S_n,$$

S_h désignant la somme

$$1 + \frac{1}{2^h} + \frac{1}{3^h} + \frac{1}{4^h} + \dots$$

A l'aide de ce résultat, j'arrive enfin au théorème suivant:

$\psi(x_1, \dots, x_n)$ étant une fonction quelconque, continue aux variables $x_1, \dots, x_n$ et satisfaisant aux conditions (A) des fonctions φ [ou aux conditions (A) et (C)], et J désignant la valeur de l'intégrale $\int \dots \int dx_1 \dots dx_n$ étendue sur le domaine $\psi(x_1, \dots, x_n) \leq 1$, on peut toujours trouver n^2 quantités réelles a_{hk} à déterminant 1, de sorte que la relation

$$0 < \psi(a_{11}y_1 + \dots + a_{1n}y_n, \dots, a_{n1}y_1 + \dots + a_{nn}y_n) \leq \sqrt[n]{\frac{S_n}{J}} \left(\text{ou} \leq \sqrt[n]{\frac{2S_n}{J}} \right)$$

ne soit vérifiée par aucun système de nombres entiers $y_1, \dots, y_n$.

En appliquant ce théorème à la fonction $\psi = \sqrt{x_1^2 + \dots + x_n^2}$, il résulte une certaine limite inférieure de la limite précise du minimum des formes quadratiques positives, et cette limite donne lieu à l'observation suivante:

$\beta_n \sqrt[n]{D}$ désignant la limite précise du minimum des formes quadratiques positives à n variables et à déterminant D , on a

$$\lim_{n \rightarrow \infty} \left(\frac{\log \beta_n}{\log n} \right) = 1.$$

XII.

Über Eigenschaften von ganzen Zahlen, die durch räumliche Anschauung erschlossen sind.

(Mathematical Papers read at the international Mathematical Congress held in connection with the world's Columbian Exposition Chicago, 1893, pp. 201—207, und, von R. Laugel ins Französische übersetzt, in den Nouvelles Annales de Mathématiques, 3^e série, t. XV, 1896 unter dem Titel: Sur les propriétés des nombres entiers qui sont dérivées de l'intuition de l'espace.)

In der Zahlentheorie wird, wie in jedem anderen Gebiete der Analysis, häufig die Erfindung mittels geometrischer Überlegungen vor sich gehen, während schließlich vielleicht nur die analytischen Verifikationen mitgeteilt werden. Ich würde deshalb schon an sich nicht in der Lage sein, mein Thema zu erschöpfen; es ist dies auch nicht meine Absicht. Ich will hier ganz allein von demjenigen geometrischen Gebilde sprechen, welches die einfachste Beziehung zu den ganzen Zahlen hat, von dem *Zahlengitter*. Darunter hat man, irgendwelche Parallelkoordinaten x, y, z im Raume vorausgesetzt, den Inbegriff derjenigen Punkte x, y, z zu verstehen, für welche x , wie y , wie z ganze Zahlen sind; der besseren Anschaulichkeit wegen denke man sich unter x, y, z gewöhnliche rechtwinklige Koordinaten.

Eine Figur, die sich als ein Ausschnitt aus dem Zahlengitter darstellt, ist es, die man beim Beweise der Multiplikationsregel $(ab)c = a(bc)$ heranzuziehen pflegt. Ich würde des weiteren die wichtigen Relationen über größte Ganze zu erwähnen haben, die Dirichlet (Crelles Journal, Bd. 47, *Über ein die Division betreffendes Problem*; Werke, Bd. II) auf geometrischem Wege erhalten hat. Ich will mich jedoch hier auf Fragen beschränken, bei denen der Begriff des Unendlichen hineinspielt, nämlich das *ganze Gitter*, nicht bloß Ausschnitte daraus in Betracht kommen. (Das Folgende gibt in der Hauptsache einiges aus meinem Buche „Geometrie der Zahlen“ (1896, bei B. G. Teubner) wieder, wobei ich bemerke, daß dort die Beschränkung auf Systeme aus *drei* ganzen Zahlen nicht statthat.)

I. Der wichtigste Begriff, der mit dem Zahlengitter in Zusammen-

hang steht, ist der des *Volumens* eines Körpers; dieser Begriff bildet dann weiter die Grundlage für den Begriff des dreifachen Integrals. Man nehme jeden Punkt des Zahlengitters zum Mittelpunkt eines Würfels mit Seitenflächen parallel den Koordinatenebenen und von der Kante 1; zu einem Würfel soll stets die Begrenzung miteingerechnet werden. Man erlangt so ein Netz N von Würfeln, welches den Raum lückenlos erfüllt, und die einzelnen Würfel darin sind untereinander in ihren inneren Punkten durchweg verschieden. Nun sei K irgendeine solche Punktmenge, welche sich ganz auf eine endliche Anzahl von Würfeln aus N verteilt. Man dilatiere diese Menge K von einem beliebigen Punkte p im Raume aus in allen Richtungen in einem beliebigen Verhältnisse $\Omega : 1$. Aus K entstehe so K_Ω^p . Sodann sei a_Ω^p die Anzahl aller der Würfel aus N , in welchen jeder einzige Punkt sich als ein *innerer* Punkt von K_Ω^p erweist, und es sei u_Ω^p die Anzahl aller Würfel aus N , welche überhaupt *mindestens einen* Punkt von K_Ω^p enthalten. Dann konvergieren nach dem, was C. Jordan (Journal de Mathématiques, 4^e série, T. 8, 1892, p. 77) gezeigt hat, immer $\Omega^{-3} \cdot a_\Omega^p$ und $\Omega^{-3} \cdot u_\Omega^p$ für ein unendlich wachsendes Ω , unabhängig von p , je nach einem bestimmten Grenzwerte A und U , dem *inneren* und dem *äußeren* Volumen von K . Man spricht vom *Volumen* von K *schlecht-hin*, wenn sich $A = U$ herausstellt.

II. Die tieferen Eigenschaften des Zahlengitters nun hängen mit einer Verallgemeinerung des Begriffs der *Länge einer geraden Linie* zusammen, bei der allein der Satz, daß in einem Dreiecke die Summe zweier Seiten niemals kleiner als die dritte ist, erhalten bleibt.

Man denke sich eine Funktion $S(ab)$ von zwei beliebig variablen Punkten a und b zunächst nur mit folgenden Eigenschaften: (1) Es soll $S(ab)$ immer positiv sein, wenn b von a verschieden ist, und Null, wenn b und a identisch sind; (2) sind a, b, c, d vier Punkte und darunter b von a verschieden, und besteht zwischen ihnen eine Beziehung $d - c = t(b - a)$ mit positivem t , so soll immer $S(cd) = tS(ab)$ sein; die genannte Beziehung ist im Sinne des baryzentrischen Kalküls aufzufassen und bedeutet, daß cd und ab Strecken von gleicher Richtung und mit Längen (im gewöhnlichen Sinne) im Verhältnisse $t : 1$ sind. Zum Unterschiede von der gewöhnlichen Länge möge $S(ab)$ *Strahldistanz von a nach b* heißen.

Es sei o der Nullpunkt; offenbar werden alle Werte $S(ab)$ festgelegt sein, sowie die Menge der Punkte u gegeben ist, für welche $S(ou) \leq 1$ ist; diese Punktmenge heiße der *Eichkörper* der Strahldistanzen, es wird zu ihm in jeder Richtung von o aus eine Strecke von o aus mit endlicher, nichtverschwindender Länge gehören müssen.

Wenn nun ferner für irgend drei Punkte a, b, c immer

$$(3) \quad S(ac) \leq S(ab) + S(bc)$$

ist, sollen die Strahldistanzen *einheitlich* heißen. Dann besitzt ihr Eichkörper die Eigenschaft, daß mit irgend zwei Punkten u, v in ihm immer die ganze Strecke uv zu diesem Körper gehört, und andererseits ist jeder *nirgends konkave Körper* mit dem Nullpunkt im Inneren Eichkörper für ganz bestimmte einheitliche Strahldistanzen.

Mit $E(ab)$ werde die halbe Kante desjenigen Würfels mit Seitenflächen parallel den Koordinatenebenen bezeichnet, der a als Mittelpunkt hat und seine Begrenzung durch b schickt. Die $E(ab)$ sind als die einfachsten einheitlichen $S(ab)$ anzusehen. Die vollständige analytische Auflösung der Bedingungen (1), (2), (3) habe ich im ersten Kapitel meiner „Geometrie der Zahlen“ gegeben. Es zeigt sich, daß auf Grund von (3) insbesondere immer die Funktion $S(ab)$ eine *stetige* der Koordinaten von a und von b ist, ferner zwei positive Größen g und G vorhanden sind, so daß man

$$gE(ab) \leq S(ab) \leq GE(ab)$$

für alle a und b hat, endlich der Eichkörper ein bestimmtes Volumen J besitzt. Die Bedeutung von g und G ist offenbar die, daß der Würfel $E(ou) \leq \frac{1}{G}$ ganz im Eichkörper enthalten ist und letzterer seinerseits ganz im Würfel $E(ou) \leq \frac{1}{g}$.

Wechselseitig sollen die $S(ab)$ heißen, wenn durchweg

$$(4) \quad S(ba) = S(ab)$$

ist. Solches hat dann und nur dann statt, wenn der Eichkörper den Nullpunkt als *Mittelpunkt* hat.

III. Es gibt im Zahlengitter offenbar Punkte r , für die $E(or) = 1$ ist. Irgendwelche *einheitliche* $S(ab)$ vorausgesetzt, wird für diese Gitterpunkte r dann $S(or) \leq G$ sein. Diese letztere Bedingung nun kann überhaupt nur von solchen Punkten r erfüllt werden, für welche $E(or) \leq \frac{G}{g}$ ist, und dieser Bedingung wieder genügen sicher nur eine endliche Anzahl Gitterpunkte. Aus diesen Gitterpunkten muß dann notwendig die *kleinste* Strahldistanz M zu ersehen sein, welche von o nach allen anderen Gitterpunkten zusammengenommen existiert und die nun jedenfalls $\leq G$ ist. Wird sodann für einen beliebigen ersten Gitterpunkt a der Körper $S(au) \leq \frac{1}{2}M$, für einen beliebigen anderen Gitterpunkt c der Körper $S(cu) \leq \frac{1}{2}M$ konstruiert, so sind solche zwei Körper zufolge (3) in ihren inneren Punkten durchweg verschieden. Werden nun die Strahldistanzen auch

noch *wechselseitig* vorausgesetzt, so ist der zweite Körper mit $S(cu) \leq \frac{1}{2} M$ identisch, und stoßen dann also die verschiedenen Körper $S(au) \leq \frac{1}{2} M$ für die verschiedenen Gitterpunkte a höchstens in den Begrenzungen zusammen.

Nun sei Ω irgendeine positive und gerade ganze Zahl, und man konstruiere die hier bezeichneten Körper für die sämtlichen im Würfel $E(ou) \leq \frac{\Omega}{2}$ enthaltenen $(\Omega + 1)^3$ Gitterpunkte

$$x, y, z = 0, \pm 1, \pm 2, \dots, \pm \frac{\Omega}{2}.$$

Aus $S(au) \leq \frac{1}{2} M \leq \frac{1}{2} G$ folgt $E(au) \leq \frac{1}{2} \frac{G}{g}$, und werden deshalb alle diese Körper in dem Würfel $E(ou) \leq \frac{1}{2} \left(\Omega + \frac{G}{g} \right)$ enthalten sein, dessen Volumen $\left(\Omega + \frac{G}{g} \right)^3$ beträgt. Indem sie nun sämtlich auseinander liegen und je vom Volumen $\left(\frac{M}{2} \right)^3 J$ sind, geht daraus die Ungleichung

$$\left(\Omega + \frac{G}{g} \right)^3 \geq (\Omega + 1)^3 \left(\frac{M}{2} \right)^3 J$$

hervor; nun stellen M und J bestimmte Größen vor und Ω kann beliebig groß genommen werden, mithin entnimmt man daraus:

$$(5) \quad 1 \geq \left(\frac{M}{2} \right)^3 J,$$

muß es also mindestens einen, von o verschiedenen Gitterpunkt q geben, für den $S(oq) \leq \frac{2}{\sqrt[3]{J}}$ ist.

Das hiermit gewonnene Theorem über die nirgends konkaven Körper mit Mittelpunkt scheint mir zu den fruchtbarsten in der ganzen Zahlentheorie zu gehören. Ich hatte es, durch das Studium der Aufsätze von Dirichlet und von Hermite über quadratische Formen (Crelles Journal, Bd. 40, S. 209 u. S. 261; Dirichlets Werke, Bd. II, S. 27; Oeuvres d'Hermite, T. I, p. 100) angeregt, zunächst für die Ellipsoide gefunden (Crelles Journal, Bd. 107, S. 291; diese Ges. Abhandlungen, Bd. I, S. 255); ein noch größeres Interesse aber bieten die Folgerungen dar, welche dieses Theorem hinsichtlich linearer Formen zuläßt und von denen ich sogleich einige hervorheben werde.

Das Gleichheitszeichen in (5) tritt dann und nur dann ein, wenn die Körper $S(au) \leq \frac{1}{2} M$ um die einzelnen Gitterpunkte a den Raum *lückenlos* erfüllen. Dazu muß vor allem die vollständige Begrenzung des Eickörpers durch eine endliche Anzahl von Ebenen, und zwar durch nicht mehr als $2(2^3 - 1)$ Ebenen, gebildet werden; nämlich es muß dann jede

ebene Wand von $S(au) \leq M$ noch exklusive des Randes mindestens einen Gitterpunkt x, y, z enthalten, und können für derartige Gitterpunkte in zwei, nicht in bezug auf o symmetrischen Wänden niemals x, y, z gleiche Reste modulo 2 ergeben, wie auch für keinen dieser Punkte $x, y, z \equiv 0, 0, 0 \pmod{2}$ sein können. Das Gleichheitszeichen in (5) tritt beispielsweise niemals für ein Oktaeder ein.

IV. Es seien ξ, η, ζ drei lineare Formen in x, y, z mit einer von Null verschiedenen Determinante D , es seien entweder alle drei reell, oder ξ reell und η, ζ zwei Formen mit konjugiert imaginären Koeffizienten; weiter sei p irgendeine reelle GröÙe. Der durch

$$(6) \quad \left(\frac{|\xi|^p + |\eta|^p + |\zeta|^p}{3} \right)^{\frac{1}{p}} \leq 1$$

definierte Körper K_p stellt dann, sowie $p \geq 1$ ist, einen nirgends konkaven Körper vor; für das Volumen J_p dieses Körpers findet man:

$$J_p = \frac{2^3}{\lambda_p^3 |D|}, \quad \lambda_p^3 = \frac{3^{-\frac{3}{p}} \Gamma\left(1 + \frac{3}{p}\right)}{\left\{ \Gamma\left(1 + \frac{1}{p}\right) \right\}^3} \quad \text{oder} \quad = \frac{2}{\pi} \frac{3^{-\frac{3}{p}} \Gamma\left(1 + \frac{3}{p}\right)}{\Gamma\left(1 + \frac{1}{p}\right) 2^{-\frac{2}{p}} \Gamma\left(1 + \frac{2}{p}\right)};$$

es zeigt sich ferner, daß für einen Körper K_p , wenn p endlich ist, in (5) niemals das Gleichheitszeichen in Betracht kommt. Man gewinnt so den Satz:

Ist $p \geq 1$, so gibt es immer ganze Zahlen x, y, z , die nicht sämtlich Null sind und für welche man

$$\left(\frac{|\xi|^p + |\eta|^p + |\zeta|^p}{3} \right)^{\frac{1}{p}} < \lambda_p |D|^{\frac{1}{3}}$$

hat.

Hält man x, y, z fest, so nimmt der Ausdruck links in (6), wenn nicht gerade $|\xi| = |\eta| = |\zeta|$ ist, in welchem Falle dieser Ausdruck von p unabhängig sein würde, mit p für alle Werte $p \geq 0$ kontinuierlich ab (sogar für alle p , wenn keine der Größen $|\xi|, |\eta|, |\zeta|$ Null ist). Es wird danach ein jeder Körper K_p in allen anderen von diesen Körpern mit kleinerem p enthalten sein und also $\frac{1}{J_p}$ und λ_p mit p kontinuierlich zunehmen; für $p = \infty$ konvergiert λ_p^3 nach 1, bzw. $\frac{2}{\pi}$. Für $p = \infty$ geht K_p in das Parallelepipedum $-1 \leq \xi \leq 1, -1 \leq \eta \leq 1, -1 \leq \zeta \leq 1$ oder den elliptischen Zylinder $-1 \leq \xi < 1, \eta^2 + \zeta^2 \leq 1$ über; K_1 hingegen stellt ein Oktaeder oder einen Doppelkegel vor. Endlich wird aus der Funktion links in (6) für $p = 0$ das geometrische Mittel $\sqrt[3]{|\xi \eta \zeta|}$, so daß man den Satz hinzufügen kann:

Es gibt immer ganze Zahlen x, y, z , die nicht sämtlich Null sind und für welche man $|\xi\eta\zeta| < \lambda_1^3 |D|$, umsomehr also $< |D|$, hat.

Diese Sätze und die analogen für n lineare Formen mit n Variablen lassen insbesondere fundamentale Anwendungen in der Theorie der algebraischen Zahlen zu, beim Beweise der Dirichletschen Sätze über die komplexen Einheiten, der Endlichkeit der Anzahl der Idealklassen, und sie haben zuerst den wichtigen Nachweis ermöglicht, daß in der Diskriminante eines jeden algebraischen Zahlkörpers immer mindestens eine Primzahl aufgeht.

V. Es seien a und b irgend zwei reelle Größen und t eine beliebige Größe > 1 . Die Anwendung der Sätze in III. auf das Parallelepipedum

$$-1 \leq x - az \leq 1, \quad -1 \leq y - bz \leq 1, \quad -1 \leq \frac{z}{t} \leq 1$$

führt dazu, daß es immer ganze Zahlen x, y, z gibt, für welche

$$0 < z \leq t^{\frac{2}{3}}, \quad |x - az| < \frac{1}{t^{\frac{1}{3}}}, \quad |y - bz| < \frac{1}{t^{\frac{1}{3}}}$$

ist. Dieses Resultat, jedoch nur für den Fall ganzzahliger Werte von t , hat bereits Kronecker (Berichte der Berliner Akademie, 1884, S. 1073; Werke, Bd. III, 1, S. 36) mittels des scheinbar trivialen, dessen ungeachtet aber äußerst erfolgreichen Prinzips (s. Dirichlet, *Verallgemeinerung eines Satzes aus der Lehre von den Kettenbrüchen*; Werke, Bd. I, S. 636) bewiesen, daß, wenn eine Anzahl von Größensystemen in eine kleinere Anzahl von Bereichen fallen, mindestens zwei Systeme darunter in einen und denselben Bereich zu liegen kommen müssen; es ist dies einer der wenigen Fälle, wo bereits dieses einfachere Prinzip wesentlich gleiche Folgerungen ermöglicht wie das arithmetische Theorem in III.

Die Betrachtung des Oktaeders

$$|x - az| + \left| \frac{z}{t} \right| \leq 1, \quad |y - bz| + \left| \frac{z}{t} \right| \leq 1$$

($t \geq 3$ vorausgesetzt), zeigt die Existenz von ganzen Zahlen x, y, z , für welche die Ausdrücke hier links beide $< \left(\frac{3}{t}\right)^{\frac{1}{3}}$ ausfallen und zugleich $z > 0$ ist, und für solche Zahlen findet man dann noch:

$$\left| \frac{x}{z} - a \right| < \frac{2}{3z^{\frac{3}{2}}}, \quad \left| \frac{y}{z} - b \right| < \frac{2}{3z^{\frac{3}{2}}}.$$

Diese Sätze weisen auf einen Weg, auf dem mit Erfolg die Ergebnisse der Lehre von den Kettenbrüchen zu verallgemeinern sind.

VI. Betrachtet man beliebige einhellige und wechselseitige $S(ab)$, so erscheint 2^3 als kleinste obere Grenze für $M^3 J$. Beschränkt man sich auf solche $S(ab)$, deren Eichkörper aus einem gegebenen Körper durch

alle möglichen linearen Transformationen hervorgehen, so findet man auch in dieser beschränkten Klasse von Funktionen bereits immer solche, für welche

$$M^3 J > 1 + \frac{1}{2^3} + \frac{1}{3^3} + \frac{1}{4^3} + \dots$$

ist. Der Nachweis dieses Satzes erfordert eine arithmetische Theorie der kontinuierlichen Gruppe aus allen linearen Transformationen.

Endlich ist zu erwähnen, daß die Ungleichung $M^3 J \leq 2^3$ für die nirgends konkaven Körper mit Mittelpunkt noch eine wesentliche Verallgemeinerung zuläßt, auf die ich indes hier nicht mehr eingehen will.

Bonn, im Juni 1893.

XIII.

Zur Theorie der Kettenbrüche.*)

(Annales de l'École Normale supérieure, 3^e série, t. XIII, pp. 41—60.)

Durch das Studium der Aufsätze von Herrn Hermite in den Bänden 40, 41 und 47 des Crelleschen Journals (Oeuvres, T. I, p. 94, p. 100, p. 164, p. 193, p. 200) bin ich zu einigen Verallgemeinerungen der Theorie der Kettenbrüche geführt, über die ich im folgenden kurz berichten will.

I. Ich werde zunächst von den Annäherungen an *eine einzelne reelle Größe* mittels rationaler Brüche sprechen.

Es sei Ω irgendein Wert ≥ 1 . Für eine beliebige reelle Größe a , welche weder eine ganze Zahl noch die Hälfte einer ganzen Zahl ist, kann man in folgender Weise eine Folge von ganzen Zahlen p_n, q_n ($n = 0, 1, 2, \dots$) bestimmen. Zuerst sei $p_0 = 1, q_0 = 0$; es sei f_0 die nächste ganze Zahl an a und $p_1 = f_0, q_1 = 1$. Sodann sei für ein $n \geq 1$ und solange als $p_n - a q_n \neq 0$ ist, ε_n das Vorzeichen von $\frac{p_{n-1} - a q_{n-1}}{p_n - a q_n}$, und es werde der absolute Betrag dieses Quotienten $= e_n + r_n$ gesetzt, so daß e_n eine ganze Zahl und $0 \leq r_n < 1$ ist; hernach mache man $s_n = e_n - \varepsilon_n \frac{q_{n-1}}{q_n}$, und, wenn $r_n = 0$ ist, $f_n = e_n$, wenn aber $r_n > 0$ ist, $f_n = e_n$ oder $= e_n + 1$, je nachdem

$$(A) \quad \frac{(s_n + 1)\Omega - 1}{1 - (1 - r_n)\Omega} \begin{matrix} > \\ \text{oder} \\ \leq \end{matrix} \frac{s_n \Omega - 1}{1 - r_n \Omega}$$

ist, endlich $p_{n+1} = f_n p_n - \varepsilon_n p_{n-1}, q_{n+1} = f_n q_n - \varepsilon_n q_{n-1}$. Man hat alsdann:

$$\frac{p_n}{q_n} = f_0 - \frac{\varepsilon_1}{f_1} \cdot \dots \cdot f_{n-2} - \frac{\varepsilon_{n-1}}{f_{n-1}}$$

und die Reihe der Zahlen p_n, q_n besitzt folgende Eigenschaften:

*) Statt der a. a. O. veröffentlichten, von L. Laugel herrührenden Übersetzung, welche den Titel trägt: *Généralisation de la théorie des fractions continues*, gelangt hier das deutsche Originalmanuskript des Verfassers zum Abdruck. (Anm. d. Herausg.)

1. Wenn a rational ist, bricht die Reihe mit irgendeinem Index ν ab, für den $\frac{p_\nu}{q_\nu} = a$ ist.

2. Man hat $0 < q_1 < q_2 < \dots$.

3. Die Zahlen p_n und q_n sind immer relativ prim; man hat

$$p_n q_{n+1} - q_n p_{n+1} = \delta_n = \varepsilon_1 \dots \varepsilon_n.$$

4. Man setze $t_0 = \infty$, und, wenn $n \geq 1$ ist,

$$|p_{n-1} - a q_{n-1}|^\Omega + t_n q_{n-1}^\Omega = |p_n - a q_n|^\Omega + t_n q_n^\Omega = T_n;$$

es sind dann $t_0, t_1, t_2, \dots$ und $T_1, T_2, \dots$ Reihen von fortwährend abnehmenden positiven Größen, und sie konvergieren nach Null, wenn a irrational ist. Dabei ist immer*)

$$T_n \leq \frac{\left[\Gamma \left(1 + \frac{2}{\Omega} \right) \right]^2}{\left[\Gamma \left(1 + \frac{1}{\Omega} \right) \right]^\Omega} \sqrt[n]{t_n}.$$

Ist a rational, so werde noch $t_{\nu+1} = 0$ gesetzt.

5. Ist t irgendein positiver Wert, der nicht in der Reihe $t_0, t_1, t_2, \dots$ vorkommt, und $t_n > t > t_{n+1}$, und ist x, y irgendein von $0, 0$, von p_n, q_n und von $-p_n, -q_n$ verschiedenes System von ganzen Zahlen, so hat man immer

$$|x - ay|^\Omega + t |y|^\Omega > |p_n - a q_n|^\Omega + t |q_n|^\Omega. **)$$

Besonders bemerkenswert sind folgende Spezialfälle dieser Entwicklung:

1) $\Omega = \infty$. Alsdann hat man für ein $n \geq 1$ immer $f_n = e_n$, $\varepsilon_{n+1} = -1$, und $\frac{p_1}{q_1}, \frac{p_2}{q_2}, \dots$ sind die Näherungsbrüche der gewöhnlichen Kettenbruchentwicklung für a , mit Ausschluß des ersten Näherungsbruches, falls der Überschuß von a über die größte in a enthaltene ganze Zahl $> \frac{1}{2}$ ist.

2) $\Omega = 2$. Die Ungleichungen (A) werden hier

$$\frac{1}{r_n} > \frac{2s_n + 1}{s_n + 2},$$

und nach der Eigenschaft 5. wird man diejenige Entwicklung in einen Kettenbruch vor sich haben, auf welche Herr Hermite im 41. Bande des Crelleschen Journals, S. 195 (Oeuvres, T. I, p. 108) geführt wurde.

*) Vgl. hierzu „Geometrie der Zahlen“ S. 122. (Anm. d. Herausg.)

**) In der französischen Übersetzung findet sich noch folgender Abschnitt:

6. Lorsque a est irrationnel et racine d'une équation du second degré à coefficients rationnels, il existe un indice l et un nombre μ tels que, à partir de $n = l$, on aura toujours

$$f_n = f_{n+\mu}, \quad \varepsilon_{n+1} = \varepsilon_{n+1+\mu}. \quad (\text{Anm. d. Herausg.})$$

3) $\Omega = 1$. Die Ungleichungen (A) gehen dann in

$$\frac{1}{r_n} + \frac{1}{s_n} \begin{matrix} > \\ \leq \end{matrix} 2$$

über. Man hat immer $T_n \leq \sqrt{2t_n}$ und tritt hier das Zeichen $=$ nur ein, wenn man $a = \frac{2PQ+1}{2Q^2}$ ($Q > 0$) hat und dabei P, Q ganze Zahlen ohne gemeinsamen Teiler sind, und zwar alsdann nur für einen Index n , für welchen $t_n = \frac{1}{2Q^2}$, $q_{n-1} < Q < q_n$ wird. Man wird danach für jeden Index n :

$$\left| \frac{p_n}{q_n} - a \right| < \frac{1}{2q_n^2}$$

haben. — Ist weiter b eine beliebige reelle Größe, so erfüllen für mindestens ein System von ganzen Zahlen X, Y , für die man

$$X - 1 < \frac{q_n b}{\delta_{n-1}} < X + 1, \quad Y - 1 < \frac{-q_{n-1} b}{\delta_{n-1}} < Y + 1$$

hat, $x = p_{n-1}X + p_n Y$, $y = q_{n-1}X + q_n Y$ die Ungleichung:

$$|x - ay - b| + t_n |y| < \sqrt{t_n}.$$

Wenn von den Gleichungen $x - ay = 0$, $x - b = 0$, $x - ay - b = 0$ keine in ganzen Zahlen x, y lösbar ist, existieren hiernach unendlich viele verschiedene ganze Zahlen x, y , für welche

$$y \neq 0, \quad |x - ay - b| < \frac{1}{4|y|}$$

ist. Dieses Theorem hat Herr Hermite im 88. Bande des Crelleschen Journals, S. 15 (Oeuvres, T. III) mit der weniger scharfen Grenze $\sqrt{\frac{2}{27}}$ an Stelle von $\frac{1}{4}$ gegeben.

II. Nunmehr betrachte ich den folgenden Ausdruck

$$\varphi = \left| \frac{\xi}{\varrho} \right|^\Omega + \left| \frac{\eta}{\sigma} \right|^\Omega + \left| \frac{\zeta}{\tau} \right|^\Omega;$$

darin seien ξ, η, ζ drei lineare Formen mit drei Variablen x, y, z , mit lauter *reellen* Koeffizienten und einer von Null verschiedenen Determinante Δ , und ϱ, σ, τ positive Parameter. Indem man anstatt ξ, η, ζ auch das System $-\xi, -\eta, -\zeta$ behandeln kann, möge $\Delta > 0$ vorausgesetzt werden. Auf diesen Ausdruck φ lassen sich dieselben Gesichtspunkte in Anwendung bringen, die mich zu den Sätzen in I. geführt haben.

Man deute x, y, z als Parallelkoordinaten (z. B. rechtwinklige) für die Punkte im Raume. Durch die Bedingung $\varphi \leq 1$ wird dann, so oft Ω einen Wert ≥ 1 hat, jedesmal ein *konvexer* Körper definiert, z. B. für

$\Omega = 1$ ein *Oktaeder*, für $\Omega = 2$ ein *Ellipsoid*, für $\Omega = \infty$ das *Parallelepipedum*, das von den sechs Ebenen

$$\xi = \pm \varrho, \quad \eta = \pm \sigma, \quad \zeta = \pm \tau$$

begrenzt wird. Dieses Parallelepipedum werde ich mit $\{\varrho, \sigma, \tau\}$ bezeichnen und seine drei begrenzenden Ebenenpaare der Reihe nach seine ξ -, η -, ζ -*Seiten* nennen.

Im vorhergehenden trat die gewöhnliche Kettenbruchentwicklung gerade bei der Annahme $\Omega = \infty$ zutage; im folgenden will ich mich deshalb auf diese Annahme beschränken.

1. Das System aller Punkte mit ganzzahligen Koordinaten x, y, z heiße das *Gitter* und ein einzelner Punkt daraus ein *Gitterpunkt*. Die Substitution $x = -x^*, y = -y^*, z = -z^*$ führt das Gitter und desgleichen ein jedes $\{\varrho, \sigma, \tau\}$ in sich über. Ein $\{\varrho, \sigma, \tau\}$ heiße *frei*, wenn es in seinem *Inneren* keinen Gitterpunkt außer dem Nullpunkte enthält. Ein freies $\{\varrho, \sigma, \tau\}$, welches diese Eigenschaft bei noch so kleiner Vergrößerung eines beliebigen seiner Parameter jedesmal verliert, werde ein *äußerstes* Parallelepipedum für ξ, η, ζ genannt; ein solches wird also mit mindestens einem Gitterpunkte im Inneren einer jeden Seitenfläche versehen sein müssen. Nach einem Satze, den ich im Bulletin des Sciences mathématiques, janvier 1893, *Lettre à M. Hermite* (diese Ges. Abhandlungen Bd. I, S. 266) zum ersten Male ausgesprochen habe, hat man in einem freien $\{\varrho, \sigma, \tau\}$ immer

$$\varrho\sigma\tau \leq \Delta,$$

d. h. ein $\{\varrho, \sigma, \tau\}$, wofür $\varrho\sigma\tau = \Delta$ ist, muß immer außer dem Nullpunkte noch weitere Gitterpunkte, sei es im Inneren, sei es nur an der Grenze, enthalten. Auf diesen Umstand kann man die Ermittlung eines äußersten $\{\varrho, \sigma, \tau\}$ gründen.

2. Eine lineare Substitution werde ich hier, ohne die Variablen zu benennen, bloß durch das quadratische System der Koeffizienten bezeichnen, wobei aus den einzelnen Gleichungen die Horizontalreihen entstammen sollen; und ein System von linearen Formen denke man sich zum Zwecke seiner Bezeichnung als eine lineare Substitution.

Um einige Besonderheiten auszuschließen, deren Berücksichtigung nur etwas mehr Raum erfordern, aber keinerlei Schwierigkeiten bereiten würde, setze ich von nun an voraus, daß von den drei Formen ξ, η, ζ keine einzige für ganzzahlige Werte x, y, z außer für 0, 0, 0 verschwinde. Es gelten dann folgende Sätze:

Ist $\{a, g, l\}$ ein äußerstes Parallelepipedum für ξ, η, ζ , so hat man immer

$$(1) \quad a g l < \Delta.$$

Es enthält $\{a, g, l\}$ genau einen Gitterpunkt auf jeder Seitenfläche, auf den gegenüberliegenden Flächen Gitterpunkte mit entgegengesetzten Koordinaten. Man kann darunter immer auf eine und nur eine Weise drei Gitterpunkte r, s, t ; r', s', t' ; r'', s'', t'' auf nicht gegenüberliegenden Flächen $\xi = \varepsilon a$, $\eta = \varepsilon' g$, $\zeta = \varepsilon'' l$ finden, so daß $\varepsilon \varepsilon' \varepsilon'' = +1$ ist, und wenn das System $\varepsilon \xi, \varepsilon' \eta, \varepsilon'' \zeta$ durch

$$P = \begin{pmatrix} r, & r', & r'' \\ s, & s', & s'' \\ t, & t', & t'' \end{pmatrix} \quad \text{in} \quad \Phi = \begin{pmatrix} a, & \pm b, & \pm c \\ \pm f, & g, & \pm h \\ \pm j, & \pm k, & l \end{pmatrix}$$

übergeht, die Größen $a, b, c, f, g, h, j, k, l$ sämtlich positiv sind und ihre Vorzeichen eines der folgenden sechs Systeme ergeben:

I.	II.	III.	IV.	V.	VI.
+ + +	+ - -	+ - -	+ + -	+ - +	+ - -
- + -	+ + +	- + -	- + +	+ + -	- + -
- - +	- - +	+ + +	+ - +	- + +	- - +

Dabei erweist sich dann die Determinante von P in den Fällen I. bis V. gleich 1, im Falle VI. gleich 0, und hat man

$$(2) \quad a > b, a > c; \quad g > h, g > f; \quad l > j, l > k,$$

und dazu noch je nach den einzelnen Fällen, die hier durch ihre Nummer kenntlich gemacht sind, folgende weitere Bedingungen:

I.	II.	III.
$b + c > a,$	$h + f > g,$	$j + k > l,$
$f > h$ oder $j > k$	$k > j$ oder $b > c$	$c > b$ oder $h > f$
IV.	V.	VI.
$b > c$ oder $h > f$ oder $j > k$	$c > b$ oder $f > h$ oder $k > j$	$b + c = a, h + f = g, j + k = l.$

Von den hier durch das Wort *oder* verbundenen zwei oder drei Bedingungen hat jedesmal *wenigstens eine* statt.

Es heiße $\{a, g, l\}$ in den Fällen I.—V. von der *ersten*, im Falle VI. von der *zweiten Art*, und von der Substitution P , welche ihrerseits $\{a, g, l\}$ vollkommen bestimmt, sage man, sie sei eine zu ξ, η, ζ gehörende Substitution.

Ist eine ganzzahlige Substitution P mit der Determinante 1 so beschaffen, daß durch sie $\varepsilon \xi, \varepsilon' \eta, \varepsilon'' \zeta$ mit geeigneten Vorzeichen $\varepsilon, \varepsilon', \varepsilon''$ in ein System Φ übergehen, welches eine der vorstehenden Bedingungen I. bis V. erfüllt, so ist sie stets eine zu ξ, η, ζ gehörende Substitution.

3. Man bilde für ein äußerstes $\{a, g, l\}$ die Systeme P und Φ . Ist darin $b > c$, so befindet sich in $\{b, g, l\}$ der Gitterpunkt r', s', t' auf dem

Rande einer ξ - und einer η -Seite und der Gitterpunkt r'', s'', t'' im Inneren einer ξ -Seite. Es muß dann $\{b, \frac{\Delta}{b\bar{l}}, l\}$, dem in (1) liegenden Satze zufolge, ein bestimmtes äußerstes $\{b, g_1, l\}$ mit einem Parameter $g_1 > g$ enthalten, und dieses wird dann *unter allen möglichen äußersten* $\{a_0, g_0, l_0\}$, *in welchen* $g_0 \geq g, l_0 \geq l$ *und* $a_0 < a$ *ist, dasjenige mit größtem Parameter* a_0 sein. Wenn $b < c$ ist, kommt die nämliche Eigenschaft einem gewissen äußersten $\{c, g, l_1\}$ ($l_1 > l$) zu. Dieses so in jedem Falle bestimmte $\{b, g_1, l\}$, beziehlich $\{c, g, l_1\}$ heiße *der* ξ -*Nachbar von* $\{a, g, l\}$. Zu diesem ξ -Nachbar kann man wieder den ξ -Nachbar bilden usf. ins Unendliche, dabei kommt man offenbar auf lauter verschiedene äußerste Parallelepipeda für ξ, η, ξ . Analog kann man sodann einen η -Nachbar und einen ξ -Nachbar von $\{a, g, l\}$ definieren. $\{a, g, l\}$ selbst wird, je nachdem $b > c$ oder $b < c$ ist, der η -Nachbar oder der ξ -Nachbar seines ξ -Nachbars sein. Nun besteht der folgende Hauptsatz:

Von einem beliebigen äußersten Parallelepipedium für ξ, η, ξ ausgehend kommt man durch fortgesetzte Bildung aller Nachbarn zu allen vorhandenen äußersten Parallelepipeda für ξ, η, ξ .

Denn es seien irgend zwei verschiedene äußerste Parallelepipeda $\{a, g, l\}$ und $\{a_0, g_0, l_0\}$ gegeben. Da keines im anderen enthalten sein kann, ist bei jedem mindestens ein Parameter größer und also auch bei einem nur *ein* Parameter größer; so sei etwa $a > a_0, g < g_0, l \leq l_0$. Die beiden Parallelepipeda haben dann $\Pi = \{a_0, g, l\}$ gemeinsam. Man bilde das System P für $\{a, g, l\}$; der Punkt r', s', t' darin kann nicht im Inneren von $\{a_0, g_0, l_0\}$ liegen und ist daher $b \geq a_0$. Wenn nun $b > c$ oder aber $b < c$ und dabei $l < l_0$ ist, wird der ξ -Nachbar von $\{a, g, l\}$, der dann $\{b, g_1, l\}$ ($g_1 > g$) oder $\{c, g, l_1\}$ ($l_1 > l$) lautet, mit $\{a_0, g_0, l_0\}$ ein Parallelepipedium Π_1 gemein haben, *welches Π enthält und dabei größer ist*. Ist aber $b < c$ und zugleich $l = l_0$, so hat der ξ -Nachbar von $\{a, g, l\}$, der dann $\{c, g, l_1\}$ ($l_1 > l$) lautet, zwar mit $\{a_0, g_0, l_0\}$ wieder nur Π , aber, indem dann $c > a_0, g < g_0, l_1 > l_0$ ist, dem eben behandelten Falle gemäß, mit dem η -Nachbar von $\{a_0, g_0, l_0\}$ gewiß ein Parallelepipedium Π_1 gemein, *das Π enthält und dabei größer ist*. Die zwei Parallelepipeda, die hier jedesmal auf das mit Π_1 bezeichnete Parallelepipedium führen, können nun identisch sein, anderenfalls operiere man mit ihnen wie mit den beiden, von welchen man ausging, usw. Dem in (1) enthaltenen Satze zufolge muß nun jedes äußerste Parallelepipedium, das Π enthält, ganz im Inneren von $\{\frac{\Delta}{g\bar{l}}, \frac{\Delta}{l\bar{a}_0}, \frac{\Delta}{a_0\bar{g}}\}$ liegen. Für die drei Parameter bei Π, Π_1 , usw. kommen daher nur eine endliche Anzahl von Werten in Frage, und eine *endliche* Anzahl von Schritten, wie der hier bezeichnete, muß zu einer

vollständigen Verbindung von $\{a, g, l\}$ und $\{a_0, g_0, l_0\}$ durch Nachbarn führen.

Es bilden so alle vorhandenen äußersten $\{a, g, l\}$ eine bestimmte *Kette, in welcher an jedem Gliede in gewisser Weise unmittelbar seine drei Nachbarn haften und dadurch ein Zusammenhang aller Glieder zustande kommt.*

Man findet die Nachbarn eines äußersten $\{a, g, l\}$ zweiter Art stets sämtlich von der ersten Art. Um die ganze zu ξ, η, ζ gehörige Kette von äußersten Parallelepipeda zu bilden, wird man nun von einem Gliede der ersten Art in ihr ausgehen; dann bedarf man nur des *Algorithmus*, durch den man von einem äußersten $\{a, g, l\}$ der ersten Art zu einem beliebigen Nachbar, und, falls dieser von der zweiten Art wird, weiter direkt zu den Nachbarn dieses Nachbars gelangt. Man bilde die Systeme P und Φ für $\{a, g, l\}$, und es sei

$$\begin{pmatrix} A, & F, & J \\ B, & G, & K \\ C, & H, & L \end{pmatrix}$$

das adjungierte System zu Φ , das System, welches symbolisch durch $\Delta\Phi^{-1}$ anzudeuten wäre. Es wird hinreichen, den ξ -Nachbar von $\{a, g, l\}$ und noch unter der Annahme $b > c$ zu behandeln, indem die übrigen möglichen Fälle aus diesem durch die geeigneten Permutationen von ξ, η, ζ und der zugehörigen Bezeichnungen hervorgehen; dabei ist zu beachten, daß in Φ den Formen ξ, η, ζ nicht bloß die Horizontalreihen, sondern ebenso die Vertikalreihen einzeln zugeordnet sind.

Man erhält, wenn $b > c$ ist, den ξ -Nachbar von $\{a, g, l\}$ durch den nachstehend dargestellten Algorithmus. Voran steht dabei jedesmal, welcher von den Fällen I. bis V. aus 2. bei dem Systeme Φ zutreffen soll. Die Klammer [] dient in der bekannten Weise als Zeichen für *größte Ganze*. Die aufgeschriebene Substitution ist jedesmal die, mit welcher P rechts zu multiplizieren ist, um die zu dem ξ -Nachbar gehörende Substitution zu erhalten; die drei Einheiten daneben sind die Quotienten aus den Einheiten $\varepsilon, \varepsilon', \varepsilon''$ für den ξ -Nachbar und für $\{a, g, l\}$; die römische Nummer darunter besagt, welche von den sechs Vorzeichenkombinationen aus 2. sich bei dem ξ -Nachbar einstellt.

Fall II. und Fall V.

Von den doppelten Vorzeichen bezieht sich das obere auf den Fall II., das untere auf den Fall V.

$$\left[\frac{G}{F}\right] = M, \quad \left[\frac{\pm H}{F}\right] = N; \quad a - bM - cN = u, \quad \pm j + kM - lN = v.$$

			m	n	δ
1)	$u < c,$	$v > k$	$M - 1$	$N + 1$	$+1$
2)	$u < b - c,$	$v < 0$	M	$N - 1$	-1
3), 4)	$u < b,$	aber nicht 1) noch 2)	M	N	
		3) $v > 0$			-1
		4) $v < 0$			$+1$
5)	$u > b,$	$v > 0$	M	$N + 1$	$+1$
6)	$u > b,$	$v < 0$	$M + 1$	N	-1
	$\begin{pmatrix} 0, & \mp \delta, & 0 \\ \pm \delta, & \mp \delta m, & 0 \\ 0, & -\delta n, & 1 \end{pmatrix}$	$\mp \delta, \mp \delta, +1;$ $\delta = +1, \text{ I}; \quad \delta = -1, \text{ IV.}$			

Fall I.

- 1) $j > k,$
- $$\begin{pmatrix} 0, & -1, & 0 \\ 1, & 1, & 0 \\ 0, & 0, & 1 \end{pmatrix} \quad +, +, +;$$
- 2) $j < k,$
- $$\begin{pmatrix} 0, & 1, & 0 \\ 1, & -1, & 0 \\ 0, & -1, & -1 \end{pmatrix} \quad +, -, -;$$

Fall III.

- 1) $a + c < 2b,$
- $$\begin{pmatrix} 0, & 1, & 0 \\ 1, & 1, & 0 \\ 0, & -1, & -1 \end{pmatrix} \quad -, +, -;$$

Fall IV.

- 1) $a < 2b, f < h, j + k < l,$
- $$\begin{pmatrix} 0, & -1, & 0 \\ 1, & 1, & 0 \\ 0, & 0, & 1 \end{pmatrix} \quad +, +, +;$$

Fall III, 2) und Fall IV., 2).

Die Bedingungen bei 1) sollen alsdann nicht erfüllt sein. Von den doppelten Vorzeichen bezieht sich im folgenden durchweg das obere auf den Fall III., das untere auf den Fall IV.

$$\begin{pmatrix} 0, & 0, & 0 \\ -1, & 1, & 0 \\ 0, & \mp 1, & \pm 1 \end{pmatrix} \quad \pm, +, \pm; \quad \text{VI.}$$

Von der zweiten Art wird also der ξ -Nachbar nur unter diesen letzten Umständen. Alsdann ist der Weg, um von $\{a, g, l\}$ direkt zu den Nachbarn des ξ -Nachbars überzugehen, durch folgende Tabelle vorgezeichnet. Darin haben jedesmal die hingeschriebene Substitution, die daneben stehenden Einheiten und römischen Ziffern dieselbe Bedeutung für das gesuchte Parallelepipedium, wie die entsprechenden Zeichen oben für den ξ -Nachbar. Die Herleitung von m, n, δ hat in den späteren Fällen nach derselben Regel wie im ersten Falle zu erfolgen.

Algorithmus für den ξ -Nachbar des ξ -Nachbars.

1) $b - c > c;$

$$\left[\frac{\pm G}{F} \right] = M, \quad \left[\frac{\pm G + H}{F} \right] = N;$$

$$a - (b - c)M - cN = u, \quad -j + (l - k)M - lN = v;$$

$$b - c = u^0, \quad c = u', \quad l - k = v'.$$

		m	n	δ
1)	$u < u', \quad v > v'$	$M - 1$	$N + 1$	$+ 1$
2)	$u < u', \quad v' > v > 0$	M	$N + 1$	$- 1$
3)	$u > u', \quad v > 0$	M	$N + 1$	$+ 1$
4)	$u < u^0, \quad v < 0$	M	N	$+ 1$
5)	$u > u^0, \quad v < 0$	$M + 1$	$N + 1$	$- 1$

$$\begin{pmatrix} 0, & \mp \delta, & 0 \\ -1, & -\delta m, & 0 \\ \pm 1, & \pm \delta(m - n), & \mp \delta \end{pmatrix} \quad \begin{matrix} \pm 1, & -\delta, & \mp \delta; \\ \delta = +1, & \text{V.}; & \delta = -1, & \text{III.} \end{matrix}$$

2) $b - c < c;$

$$\left[\frac{\pm K + L}{J} \right] = M, \quad \left[\frac{\pm K}{L} \right] = N;$$

$$a - cM - (b - c)N = u, \quad \pm f + hM - (g + h)N = v;$$

$$c = u^0, \quad b - c = u', \quad h = v'.$$

$$\begin{pmatrix} 0, & 0, & \mp \delta \\ 0, & -\delta, & -\delta n \\ \mp 1, & \pm \delta, & \mp \delta(m - n) \end{pmatrix} \quad \begin{matrix} \pm 1, & -\delta, & \mp \delta; \\ \delta = +1, & \text{IV.}; & \delta = -1, & \text{II.} \end{matrix}$$

Der η -Nachbar des ξ -Nachbars ist $\{a, g, l\}$ selbst.

Algorithmus für den ξ -Nachbar des ξ -Nachbars.

1) $k < l - k;$

$$\left[\frac{-H}{F} \right] = M, \quad \left[\frac{\mp G - H}{F} \right] = N;$$

$$j - (l - k)M - kN = u, \quad -a + (b - c)M - bN = v;$$

$$l - k = u^0, \quad k = u', \quad b - c = v'.$$

$$\begin{pmatrix} 0, & \mp \delta, & 0 \\ \delta, & -\delta(m - n), & -1 \\ 0, & \pm \delta m, & \pm 1 \end{pmatrix} \quad \begin{matrix} \mp \delta, & -\delta, & \pm 1; \\ \delta = +1, & \text{IV.}; & \delta = -1, & \text{I.} \end{matrix}$$

2) $k > l - k;$

2)₁ $j > k,$

$$\begin{pmatrix} \pm 1, & 0, & 0 \\ 0, & 1, & 1 \\ \mp 1, & \mp 1, & 0 \end{pmatrix} \quad \begin{matrix} \pm, & +, & \pm; \\ & \text{II.} & \end{matrix}$$

2)₂ $j < k,$

$$\text{im Falle III., 2):} \quad \begin{pmatrix} -1, & 0, & 0 \\ -1, & -1, & 1 \\ 1, & 1, & 0 \end{pmatrix} \quad \begin{matrix} -, & -, & +; \\ & \text{V.} & \end{matrix}$$

$$\text{im Falle IV., 2):} \quad \begin{pmatrix} 1, & 0, & 0 \\ 0, & -1, & 1 \\ 0, & -1, & 0 \end{pmatrix} \quad \begin{matrix} +, & -, & -; \\ & \text{V.} & \end{matrix}$$

4. Es sei ein reeller algebraischer Zahlkörper dritten Grades Θ gegeben, dessen konjugierte Körper Θ' , Θ'' ebenfalls reell sind, also mit einer *positiven* Diskriminante D . Es seien α, β, γ drei ganze Zahlen aus Θ von solcher Art, daß die Form $\xi = \alpha x + \beta y + \gamma z$ für die rationalen ganzzahligen Werte von x, y, z *alle* ganzen Zahlen aus Θ darstellt; $\eta = \xi'$ und $\zeta = \xi''$ seien die konjugierten Formen zu ξ in den Körpern Θ' und Θ'' . Die Determinante von ξ, η, ζ ist dann $\sqrt{D}$, und zwar möge sie gleich dem positiven Werte dieser Wurzel angenommen werden, indem auch $-\alpha, -\beta, -\gamma$ die Stelle von α, β, γ übernehmen können. Das Produkt $\xi\eta\zeta = Nm\xi$ ist eine Form in x, y, z mit lauter *rationalen ganzzahligen* Koeffizienten von der Diskriminante D . Wendet man auf diese Form eine Substitution P der zu ξ, η, ζ gehörigen Kette an, so entsteht eine Form φ , wieder mit rationalen ganzzahligen Koeffizienten, von der Diskriminante D oder 0, je nachdem die Determinante von P Eins oder Null ist; dabei ergeben sich aus den Ungleichungen 2.(1) und 2.(2) gewisse, nur von D abhängige obere Grenzen für die Beträge aller Koeffizienten in φ . *Es gehen danach aus $Nm\xi$ durch die sämtlichen unendlich vielen Substitutionen P überhaupt nur eine endliche Anzahl verschiedener Formen*

φ hervor. Unter einer *Einheit* des Körpers Θ soll eine ganze Zahl aus Θ mit der Norm 1 verstanden werden.

Es seien nun P und Q zwei verschiedene Substitutionen der Kette zu ξ, η, ζ , welche $\xi\eta\zeta$ in ein und dieselbe Form φ transformieren. Durch P mögen ξ, η, ζ in Ξ, H, Z übergehen; durch Q müssen dann ξ, η, ζ in dieselben Formen bis auf Faktoren übergehen, und dabei kann mit Rücksicht auf die Ungleichungen 2.(2) auch keine Änderung in der Reihenfolge der Formen eintreten, so daß aus ξ, η, ζ durch Q der Reihe nach wird $\omega\Xi, \omega'H, \omega''Z$. Dabei erweisen sich dann $\omega, \omega', \omega''$ als konjugierte Zahlen aus den Körpern $\Theta, \Theta', \Theta''$ und hat man $\omega\omega'\omega'' = 1$. Der Faktor ω wird also eine Einheit darstellen, sowie er eine ganze algebraische Zahl ist. Dies wird nun immer der Fall sein, wenn P und Q die Determinante 1 haben. Denn alsdann geht durch QP^{-1} das System ξ, η, ζ in $\omega\xi, \omega'\eta, \omega''\zeta$, und also, wenn E die identische Substitution, w einen unbestimmten Parameter bedeutet, durch $QP^{-1} - wE$ (bei Anwendung einer bekannten Symbolik) ξ, η, ζ in $(\omega - w)\xi, (\omega' - w)\eta, (\omega'' - w)\zeta$ über; danach ist die Determinante von $QP^{-1} - wE$ gleich $(\omega - w)(\omega' - w)(\omega'' - w)$ und diese Relation erweist ω als ganze Zahl.

Auf zwei verschiedene Substitutionen P und Q von der Determinante 1, welche $Nm\xi$ in ein und dieselbe Form φ transformieren, wird man z. B. mit Hilfe einer hinreichend verlängerten solchen Reihe von äußersten Parallelepipeda kommen können, in welcher jedes Parallelepipedium der ξ -Nachbar des vorhergehenden ist. Dabei stellt sich dann offenbar eine Einheit ω heraus, für welche von den Beträgen der Zahlen $\omega, \omega', \omega''$ der erste < 1 , der zweite und dritte > 1 sind. Analog kann man eine Einheit ω finden, für welche von diesen Beträgen der zweite < 1 , der dritte und erste > 1 , oder endlich der dritte < 1 , der erste und zweite > 1 sind. Es leuchtet ein, daß von solchen drei Einheiten je zwei immer unabhängig, d. h. nicht als Potenzen einer einzigen Einheit darstellbar sind.

Durch eine Substitution O von der Determinante 1 geht das Gitter immer in sich selbst über, und wird aus einem äußersten Parallelepipedium mit einer Substitution P daher wieder ein äußerstes Parallelepipedium, mit der Substitution $O^{-1}P$, das freilich nicht derselben Kette anzugehören braucht. Ist andererseits ω eine ganze Zahl aus Θ , so geht immer $\omega\xi$ aus ξ durch eine bestimmte ganzzahlige Substitution O hervor; durch dieselbe geht dann $\omega'\eta$ aus η und $\omega''\zeta$ aus ζ hervor, und wird dabei die Determinante von O also gleich $\omega\omega'\omega''$, d. i. gleich 1, sowie ω eine Einheit vorstellt. Daraus ersieht man nun: Ist $\{\lambda, \mu, \nu\}$ irgendein äußerstes Parallelepipedium für ξ, η, ζ , P die dazu gehörige Substitution, ω irgendeine Einheit aus Θ , O die ganzzahlige Substitution, welche ξ in $\omega\xi$ transformiert,

so ist auch $\left\{ \frac{\lambda}{|\omega|}, \frac{\mu}{|\omega'|}, \frac{\nu}{|\omega''|} \right\}$ immer ein äußerstes Parallelepipedum für ξ, η, ζ , es gehört dazu die Substitution $O^{-1}P$ und transformiert diese $Nm\xi$ in genau dieselbe Form φ wie P . Zwei in der hier erörterten Beziehung zueinander stehende äußerste Parallelepipeda mögen äquivalent heißen.

Es entspringt daraus nun folgendes Verfahren, alle Einheiten des Körpers Θ zu finden. Man gehe von irgendeinem äußersten Parallelepipedum für ξ, η, ζ aus, bilde die Form φ dazu, konstruiere einen Nachbar, bilde die Form φ für ihn, und man setze immer von den erhaltenen Parallelepipeda aus die Bildung von Nachbarn und der Formen φ dazu fort, soweit als dies angeht, ohne daß man zwei Parallelepipeda erster Art mit derselben Form φ und zudem beide mit zwei gleichbenannten Nachbarn in der Reihe hat. Man kommt so notwendig auf eine begrenzte Anzahl von äußersten Parallelepipeda, welche man eine *Fundamentalreihe* für die zu ξ, η, ζ gehörige Kette nennen kann. Man bilde nun für zwei unter den erhaltenen Parallelepipeda $\{\lambda, \mu, \nu\}$, welchen dieselbe Form φ entspricht, immer den Quotienten aus ihren Parametern λ ; falls dieser eine ganze Zahl wird, stellt er, mit einem geeigneten Vorzeichen versehen, eine Einheit vor; man kommt so auf eine endliche Anzahl von Einheiten, aus welchen durch Multiplikation und Division alle vorhandenen Einheiten des Körpers Θ abzuleiten sind. Es müssen sich nach dem Obigen darunter gewiß zwei unabhängige Einheiten finden, und kann man immer leicht auch zwei solche Einheiten ermitteln, aus welchen allein schon durch Multiplikation und Division *alle* Einheiten hervorgehen. —

Der in diesem Aufsätze dargelegte Algorithmus, um die Einheiten in einem kubischen Körper mit positiver Diskriminante zu finden, ist durchaus analog der Lösung der Pellschen Gleichung durch Bildung einer Periode reduzierter indefiniter binärer quadratischer Formen, wofern man die Reduktionsbedingungen von Gauß verwendet. Auf die kubischen Körper mit negativer Diskriminante kann ich an dieser Stelle nicht mehr eingehen.

Der einfachste Körper Θ wird durch $2\cos\frac{2\pi}{7}$ bestimmt*); $\vartheta = 2\cos\frac{2\pi}{7}$, $\vartheta' = 2\cos\frac{4\pi}{7}$, $\vartheta'' = 2\cos\frac{6\pi}{7}$ sind drei konjugierte ganze algebraische Zahlen; sie haben angenähert die Werte

$$\vartheta = 1,25, \quad \vartheta' = -0,45, \quad \vartheta'' = -1,80$$

*) In der Abhandlung *De la réduction des formes quadratiques ternaires positives et de son application aux irrationnelles du troisième degré* (Annales de l'École Normale supérieure, 2^e série, supplément au tome IX) hat Herr L. Charve unter anderen Beispielen diesen Körper ebenfalls behandelt.

und sind danach die Wurzeln der Gleichung

$$\vartheta^3 + \vartheta^2 - 2\vartheta - 1 = 0.$$

Die Diskriminante dieser Gleichung ist 49, also ein Quadrat, Θ somit ein Abelscher Körper. Man hat

$$(\vartheta' - \vartheta)(\vartheta'' - \vartheta)(\vartheta'' - \vartheta') = -7,$$

und entnimmt daraus

$$\vartheta' = \vartheta^2 - 2, \quad \vartheta'' = -\vartheta^2 - \vartheta + 1$$

und die aus diesen durch zyklische Permutation von $\vartheta, \vartheta', \vartheta''$ hervorgehenden Relationen. Man kann nun oben

$$\alpha, \beta, \gamma = -1, \quad -\vartheta, \quad -\vartheta^2$$

nehmen. Dann sind die Koeffizienten in ξ, η, ζ :

$$\begin{pmatrix} -1, & -1,25, & -1,55 \\ -1, & 0,45, & -0,20 \\ -1, & 1,80, & -3,25 \end{pmatrix}.$$

Es gehen jetzt $-\xi, -\eta, \zeta$ durch

$$P = \begin{pmatrix} 0, & 1, & 1 \\ 1, & 0, & 0 \\ 0, & 0, & -1 \end{pmatrix}$$

in

$$\Phi = \begin{pmatrix} 1,25, & 1, & -0,55 \\ -0,45, & 1, & 0,80 \\ 1,80, & -1, & 2,25 \end{pmatrix} = \begin{pmatrix} \vartheta, & 1, & 1 - \vartheta^2 \\ \vartheta', & 1, & 1 - \vartheta'^2 \\ -\vartheta'', & -1, & -1 + \vartheta''^2 \end{pmatrix}$$

über. Dieses System Φ genügt den Bedingungen 2. IV. und ist somit $\{\vartheta, 1, -1 + \vartheta''^2\} = (\Phi)$ ein äußerstes Parallelepipedium zu ξ, η, ζ und P die dazu gehörige Substitution. Es geht sodann $\xi\eta\zeta$ durch P in

$$\varphi = -x^3 - y^3 - z^3 + 2x^2y + 2y^2z + 2z^2x + xy^2 + yz^2 + zx^2 + xyz$$

über. Der Umstand, daß diese Form φ bei den zyklischen Permutationen von x, y, z ungeändert bleibt, setzt in Evidenz, daß Θ ein Abelscher Körper ist. *Man nimmt hier auch sofort eine charakteristische Eigenschaft aller in analoger Weise in bezug auf Abelsche Körper gebildeten Ketten wahr.*

Bei der Bestimmung des ξ -, wie des η -, wie des ζ -Nachbars von (Φ) wird nun jedesmal die gleiche Regel Anwendung finden, hier die Regel für den Fall IV., 2), so daß diese Nachbarn sämtlich von der zweiten Art werden. Dabei wird aus Φ :

$$\begin{aligned}\Psi &= \begin{pmatrix} 1, & -0,45, & -0,55 \\ -1, & 1,80, & -0,80 \\ -1, & -1,25, & 2,25 \end{pmatrix} \\ \Psi' &= \begin{pmatrix} 1,25, & -0,55, & -0,70 \\ -0,45, & 0,80, & -0,35 \\ -1,80, & -2,25, & 4,05 \end{pmatrix} \\ \Psi'' &= \begin{pmatrix} 2,25, & -1, & -1,25 \\ -0,55, & 1, & -0,45 \\ -0,80, & -1, & 1,80 \end{pmatrix},\end{aligned}$$

und φ geht jedesmal in dieselbe Form

$$\psi = x^3 + y^3 + z^3 - x^2y - y^2z - z^2x - 2xy^2 - 2yz^2 - 2zx^2 + 2xyz$$

von der Diskriminante Null über. Der ξ -Parameter in den zugehörigen äußersten Parallelepipeda (Ψ) , (Ψ') , (Ψ'') ist 1 , ϑ , $1 + \vartheta$, und indem die Quotienten dieser Größen sich als ganze Zahlen erweisen, sind diese Parallelepipeda äquivalent. Nun ist umgekehrt (Φ) der η -, ξ -, ξ -Nachbar von (Ψ) , (Ψ') , (Ψ'') und werden daher alle Nachbarn von (Ψ) , (Ψ') , (Ψ'') mit (Φ) äquivalente Parallelepipeda sein. In (Φ) , (Ψ) , (Ψ') , (Ψ'') hat man somit bereits eine Fundamentalreihe der zu ξ , η , ξ gehörigen Kette erlangt, und man kommt zu dem Resultate, daß die Einheiten ϑ , $\frac{-1}{1+\vartheta} = \vartheta'$, $\frac{-1-\vartheta}{\vartheta} = \vartheta''$, zwischen denen noch die Beziehung $\vartheta\vartheta'\vartheta'' = 1$ besteht, durch ihre Potenzen und deren Produkte alle Einheiten des Körpers Θ ergeben. Es entsprechen diesen Einheiten die Transformationen

$$\begin{pmatrix} 0, & 1, & -1 \\ 1, & 0, & 0 \\ -1, & 0, & -1 \end{pmatrix}, \quad \begin{pmatrix} 0, & 0, & 1 \\ 0, & -1, & -1 \\ 1, & -1, & 0 \end{pmatrix}, \quad \begin{pmatrix} -1, & -1, & 0 \\ -1, & 0, & 1 \\ 0, & 1, & 0 \end{pmatrix}$$

der Form φ in sich selbst. —

Mit leichten Modifikationen lassen sich die Sätze der Abschnitte 2. und 3. auch auf solche lineare Formen ausdehnen, durch welche die Null rational darstellbar ist. Für drei Formen von der besonderen Gestalt

$$x + ay + bz, \quad y + cz, \quad z$$

hat man offenbar immer in $\{1, 1, 1\}$ ein äußerstes Parallelepipedum und damit einen ganz bestimmten Ausgangspunkt für die zu den Formen gehörige Kette. Aus den Sätzen dieses Abschnittes entnimmt man leicht, daß, wenn α , β , γ , ξ , η , ζ die oben festgesetzte Bedeutung für den kubischen Körper Θ haben, jede Substitution P der zu ξ , η , ζ gehörenden Kette, in deren Parallelepipedum $\{\lambda, \mu, \nu\}$ die Quotienten $\frac{\mu}{\lambda}$, $\frac{\nu}{\mu}$ gewisse Größen

übersteigen, auch in der zu den Formen

$$x + \frac{\beta}{\alpha}y + \frac{\gamma}{\alpha}z, \quad y + \frac{\alpha\gamma' - \alpha'\gamma}{\alpha\beta' - \alpha'\beta}z, \quad z$$

gehörigen Kette auftreten muß. Derjenige Satz, welcher diesem bei zwei linearen Formen entspricht, ist genau der Satz von Lagrange, daß für eine reelle quadratische Irrationalzahl die Entwicklung in einen gewöhnlichen Kettenbruch sich periodisch gestaltet.

Ich werde bei nächster Gelegenheit auf die Untersuchung dreier Formen von der Gestalt

$$x - az, \quad y - bz, \quad z,$$

worin a, b beliebige reelle Größen sind, und eine andere, damit zusammenhängende und weit bemerkenswertere Verallgemeinerung dieses Satzes von Lagrange zurückkommen.

Königsberg, den 15. Oktober 1894.

XIV.

Ein Kriterium für die algebraischen Zahlen.

(Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen.
Mathematisch-physikalische Klasse. 1899. S. 64—88.)

(Vorgelegt in der Sitzung vom 11. Februar 1899 von D. Hilbert.)

Im Jahre 1770 hat Lagrange*) gezeigt, daß die Entwicklung einer reellen irrationalen Größe in einen gewöhnlichen Kettenbruch immer dann und nur dann periodisch ausfällt, wenn die Größe Wurzel einer quadratischen Gleichung mit rationalen Koeffizienten ist. Dieser Satz gibt offenbar ein vollständiges Mittel zur Unterscheidung der reellen algebraischen Zahlen zweiten Grades von allen anderen Größen. Seit jener Entdeckung von Lagrange durfte man vermuten, daß ein allgemeinerer Satz existiere, der ein vollständiges Kriterium für die reellen (oder komplexen) algebraischen Zahlen beliebigen n^{ten} Grades gibt und der für $n = 2$ und reelle Zahlen eben auf jenen Satz von Lagrange hinauskommt. Eine solche Verallgemeinerung wird zum ersten Male**) im folgenden dargelegt.

§ 1. Arithmetische Hilfssätze.

1. Ich beginne mit der Ableitung einiger Hilfssätze, auf welche sich die späteren Beweisführungen gründen werden.

Es seien $\xi_1, \dots, \xi_v$ eine Reihe linearer homogener Formen mit den n reellen Variablen $x_1, \dots, x_n$ und mit irgendwelchen reellen oder komplexen Koeffizienten; nur soll das System der Gleichungen $\xi_1 = 0, \dots, \xi_v = 0$ bloß durch das *eine reelle* Wertsystem $x_1 = 0, \dots, x_n = 0$ befriedigt werden können.

Der größte unter den absoluten Beträgen von $x_1, \dots, x_n$ vorkommende Betrag soll mit $\max |x_k|$ bezeichnet werden. Setzt man für $x_1, \dots, x_n$ irgendwelche reellen Werte, so soll der größte unter den absoluten Be-

*) Abhandlungen der Akademie zu Berlin, Bd. XXIV, 1770; Werke, Bd. II, S. 603 ff.

**) Über bisherige Versuche in dieser Richtung s. P. Bachmann, Vorlesungen über die Natur der Irrationalzahlen, 1892, Vorl. II und Vorl. X.

trägen von $\xi_1, \dots, \xi_v$ vorkommende Betrag mit $\max |\xi_x(x_1, \dots, x_n)|$ und zugleich mit $f(x_1, \dots, x_n)$ bezeichnet werden.

Man hat dann

$$(1) \quad f(-x_1, \dots, -x_n) = f(x_1, \dots, x_n),$$

$$(2) \quad f(tx_1, \dots, tx_n) = tf(x_1, \dots, x_n), \text{ wenn } t > 0 \text{ ist.}$$

Die Funktion $f(x_1, \dots, x_n)$ ist eine stetige der Argumente $x_1, \dots, x_n$ und hat daher in dem durch $\max |x_k| = 1$ definierten *abgeschlossenen* Bereiche (d. i. auf der *Begrenzung* des durch $-1 \leq x_1 \leq 1, \dots, -1 \leq x_n \leq 1$ definierten Würfels) ein bestimmtes Minimum g , das wegen der an die $\xi_1, \dots, \xi_v$ oben gestellten Anforderung gewiß > 0 ist, und ein bestimmtes Maximum G . Sodann ist wegen (2) stets

$$(3) \quad g \max |x_k| \leq f(x_1, \dots, x_n) \leq G \max |x_k|.$$

Endlich hat man, wenn $a_1, \dots, a_n$ und $b_1, \dots, b_n$ zwei reelle Systeme sind, stets

$$(4) \quad f(a_1 + b_1, \dots, a_n + b_n) \leq f(a_1, \dots, a_n) + f(b_1, \dots, b_n).$$

Denn für jede einzelne der linearen Formen ξ_i gilt

$$\begin{aligned} |\xi_i(a_1 + b_1, \dots, a_n + b_n)| &\leq |\xi_i(a_1, \dots, a_n)| + |\xi_i(b_1, \dots, b_n)| \\ &\leq \max |\xi_x(a_1, \dots, a_n)| + \max |\xi_x(b_1, \dots, b_n)| \end{aligned}$$

und daher auch

$$\max |\xi_x(a_1 + b_1, \dots, a_n + b_n)| \leq \max |\xi_x(a_1, \dots, a_n)| + \max |\xi_x(b_1, \dots, b_n)|.$$

Hat man $f(a_1, \dots, a_n) \leq 1$ und $f(b_1, \dots, b_n) \leq 1$ und ist $0 < t < 1$, so folgt nach den Regeln (4) und (2)

$$(5) \quad f((1-t)a_1 + tb_1, \dots, (1-t)a_n + tb_n) \leq (1-t)f(a_1, \dots, a_n) + tf(b_1, \dots, b_n) \leq 1.$$

Nach den Eigenschaften (5) und (1) ist der durch $f(x_1, \dots, x_n) \leq 1$ definierte Bereich K in der Mannigfaltigkeit der $x_1, \dots, x_n$ ein *nirgends konkaver Körper* und hat das System $x_1 = 0, \dots, x_n = 0$ (den Nullpunkt) als *Mittelpunkt*. Der Bereich K liegt wegen (3) ganz im Würfel $\max |x_k| \leq \frac{1}{g}$ eingeschlossen und enthält in sich den Würfel $\max |x_k| \leq \frac{1}{G}$. Das n -fache Integral $\int dx_1 \dots dx_n$, über den Bereich K erstreckt, (das Volumen von K) hat einen bestimmten positiven endlichen Wert J .*)

2. Der Inbegriff aller Systeme $x_1, \dots, x_n$, bei welchen sowohl x_1 , wie $x_2, \dots$, wie x_n ganze Zahlen sind, soll das *Zahlengitter*, die einzelnen Systeme daraus sollen *Gitterpunkte* heißen.

*) Man kann die Zahl v oben auch unbegrenzt wachsen lassen, wenn man die Bedingung hinzufügt, daß in allen Formen ξ_i die Beträge der Koeffizienten unter einer Grenze bleiben, und kommt dadurch zu dem Begriffe eines beliebigen nirgends konkaven Körpers mit dem Nullpunkt als Mittelpunkt. Alle im § 1 abgeleiteten Sätze gelten unverändert für jeden solchen Körper.

Eine Reihe von Systemen $x_1 = p_1^{(h)}, \dots, x_n = p_n^{(h)}$ ($h = 1, \dots, m$ und $m \leq n$) soll *unabhängig* heißen, wenn in der aus ihnen zu bildenden Matrix $\|p_k^{(h)}\|$ nicht jede m -reihige Determinante Null ist.

Es seien $p_1^{(h)}, \dots, p_n^{(h)}$ für $h = 1, \dots, n$ irgend n unabhängige Gitterpunkte, also die Substitution P :

$$(6) \quad x_k = p_k^{(1)} z_1 + \dots + p_k^{(n)} z_n \quad (k = 1, \dots, n)$$

eine ganzzahlige mit von Null verschiedener Determinante. Dann gibt es bekanntlich eine ganzzahlige Substitution A mit einer Determinante $= \pm 1$:

$$(7) \quad x_k = a_k^{(1)} y_1 + \dots + a_k^{(n)} y_n \quad (k = 1, \dots, n)$$

so daß die Formeln $P^{-1}A$ werden:

$$(8) \quad z_1 = \gamma_1^{(1)} y_1 + \dots + \gamma_1^{(n)} y_n, \dots, z_n = \gamma_n^{(n)} y_n \quad (\gamma_h^{(k)} = 0, h > k)$$

und dabei ferner die Ungleichungen erfüllt sind:

$$(9) \quad 0 < \gamma_h^{(h)} \leq 1, \quad 0 \leq \gamma_h^{(k)} < \gamma_h^{(h)} (h < k).$$

In der Tat, unter allen Gitterpunkten, deren Koordinaten von der Form $x_k = \gamma_1 p_k^{(1)}$ mit $0 < \gamma_1 \leq 1$ ($k = 1, \dots, n$) sind, wird es einen geben, für den γ_1 am kleinsten ist; es sei für ihn $\gamma_1 = \gamma_1^{(1)}$, $x_k = a_k^{(1)}$. Dann kann man unter allen Gitterpunkten mit Koordinaten von der Form $x_k = \gamma_1 p_k^{(1)} + \gamma_2 p_k^{(2)}$ ($k = 1, \dots, n$) und den Umständen $0 \leq \gamma_1 < \gamma_1^{(1)}$, $0 < \gamma_2 \leq 1$ einen finden, für den γ_2 so klein als möglich ist; es sei für ihn $\gamma_2 = \gamma_2^{(2)}$, $\gamma_1 = \gamma_1^{(2)}$, $x_k = a_k^{(2)}$, usf. Zuletzt kann man unter den Gitterpunkten mit Koordinaten von der Form $x_k = \gamma_1 p_k^{(1)} + \dots + \gamma_n p_k^{(n)}$ und $0 \leq \gamma_1^{(1)}, \dots, 0 \leq \gamma_{n-1} < \gamma_{n-1}^{(n-1)}$, $0 < \gamma_n \leq 1$ einen finden, für den γ_n möglichst klein ist; für ihn sei $\gamma_n = \gamma_n^{(n)}, \dots, \gamma_1 = \gamma_1^{(n)}$, $x_k = a_k^{(n)}$.

Nunmehr wird man, wenn x_k ($k = 1, \dots, n$) ein beliebiger Gitterpunkt ist, sukzessive $y_n, y_{n-1}, \dots, y_1$ als ganze Zahlen so bestimmen können, daß in der Umformung

$$x_k - (y_n a_k^{(n)} + y_{n-1} a_k^{(n-1)} + \dots + y_1 a_k^{(1)}) = \gamma_n p_k^{(n)} + \gamma_{n-1} p_k^{(n-1)} + \dots + \gamma_1 p_k^{(1)} \\ (k = 1, \dots, n)$$

sich $0 \leq \gamma_n < \gamma_n^{(n)}$, $0 \leq \gamma_{n-1} < \gamma_{n-1}^{(n-1)}, \dots, 0 \leq \gamma < \gamma_1^{(1)}$ ergibt. Dann muß, da die linken Seiten jedenfalls Koordinaten eines Gitterpunktes sind, nach der Bedeutung von $\gamma_n^{(n)}, \dots, \gamma_1^{(1)}$ hier notwendig $\gamma_n = 0$, $\gamma_{n-1} = 0, \dots, \gamma_1 = 0$ sein; es folgen also die Relationen von der Form (7), woraus nach den Ausdrücken der $a_k^{(h)}$ weiter die Formeln (8) hervorgehen. Zu jedem Systeme von ganzen Zahlen $x_1, \dots, x_n$ gehören so vermöge (7) bestimmte ganzzahlige Werte $y_1, \dots, y_n$. Es müssen also in der Auflösung von (7):

$$(10) \quad y_h = b_h^{(1)} x_1 + \dots + b_h^{(n)} x_n \quad (h = 1, \dots, n),$$

wie die Einführung der n Systeme $x_j = 1, x_k = 0$ ($k \neq j$) für $j = 1, \dots, n$ zeigt, alle Koeffizienten $b_k^{(j)}$ ganze Zahlen sein. Somit ist auch ihre Determinante $|b_k^{(j)}|$ eine ganze Zahl, und da dieselbe der reziproke Wert der ebenfalls ganzzahligen Determinante $|a_k^{(h)}|$ ist, so folgt notwendig, daß diese: $|a_k^{(h)}| = \pm 1$ ist.

3. Betrachten wir wieder die in 1. definierte Funktion $f = f(x_1, \dots, x_n)$. Es gibt wegen (3) nur eine endliche Anzahl von Gitterpunkten, für welche f eine gegebene Größe nicht überschreitet. Andererseits gilt nach (3) $f \leq G$ jedenfalls bei den n unabhängigen Systemen $x_j = 1, x_k = 0$ ($k \neq j$) für $j = 1, \dots, n$. Man bestimme nun unter allen vom Nullpunkte verschiedenen Gitterpunkten, bei welchen $f \leq G$ ist, einen ersten Gitterpunkt $p_1^{(1)}, \dots, p_n^{(1)}$, so daß $f(p_1^{(1)}, \dots, p_n^{(1)}) = F_1$ möglichst klein ist, sodann einen zweiten, von diesem ersten unabhängigen Gitterpunkt $p_1^{(2)}, \dots, p_n^{(2)}$, so daß $f(p_1^{(2)}, \dots, p_n^{(2)}) = F_2$ möglichst klein ist, usf. bis zu einem n^{ten} Gitterpunkt $p_1^{(n)}, \dots, p_n^{(n)}$, so daß schließlich die Determinante $|p_k^{(h)}| \neq 0$ ist und für diesen letzten $f(p_1^{(n)}, \dots, p_n^{(n)}) = F_n$ möglichst klein ausfällt. Bei jedem einzelnen der zu wählenden Gitterpunkte hat man jedenfalls die Auswahl zwischen einem Paare entgegengesetzter Systeme $(p_1, \dots, p_n$ und $-p_1, \dots, -p_n)$, unter Umständen aber zwischen einer gewissen Anzahl solcher Paare; trotz der dabei zugelassenen Willkür aber ist das System der n Werte $F_1, F_2, \dots, F_n$ von vornherein ein völlig bestimmtes.

In der Tat, jedenfalls ist

$$(11) \quad F_1 \leq F_2 \leq \dots \leq F_n.$$

Wir denken uns nun für die n Gitterpunkte $p_1^{(h)}, \dots, p_n^{(h)}$ ($h = 1, \dots, n$), an welche die Betrachtungen in 2. anknüpfen, die hier ausgewählten n Gitterpunkte gesetzt und können alsdann sämtliche dort eingeführten Bezeichnungen hier übernehmen. Vermöge der Substitution A entsprechen sich genau die Gitterpunkte $x_1, \dots, x_n$ und die ganzzahligen Systeme $y_1, \dots, y_n$. Nach der Bedeutung der Werte F_j hat man für einen Gitterpunkt, bei dem $z_j, z_{j+1}, \dots, z_n$ nicht sämtlich Null sind oder also $y_j, y_{j+1}, \dots, y_n$ nicht sämtlich Null sind, stets $f(x_1, \dots, x_n) \geq F_j$. Hat man nun irgend n unabhängige Gitterpunkte $x_1^{(h)}, \dots, x_n^{(h)}$ für $h = 1, \dots, n$, so daß also die Determinante $|x_k^{(h)}| \neq 0$ ist, und entsprechen ihnen vermöge der Substitution (7) die Systeme $y_1^{(h)}, \dots, y_n^{(h)}$, so ist auch die Determinante $|y_k^{(h)}| \neq 0$; es können daher, wenn j einen der Werte $1, \dots, n$ bedeutet, nicht bei j oder gar mehr der Punkte alle $n - j + 1$ Größen $y_j, y_{j+1}, \dots, y_n$ gleich Null sein, es sind also stets für mindestens $n - j + 1$ der Punkte die Werte $f(x_1, \dots, x_n) \geq F_j$. Ordnet man also die n Werte $f(x_1^{(h)}, \dots, x_n^{(h)})$ der Größe nach in die Reihe $F_1^*, \dots, F_n^*$, so ist stets

$F_j^* \geq F_j$ ($j = 1, \dots, n$), [also auch $\prod_{h=1}^n f(x_1^{(h)}, \dots, x_n^{(h)}) \geq F_1 \dots F_n$, wobei das Gleichheitszeichen nur statthat, wenn die n Werte $f(x_1^{(h)}, \dots, x_n^{(h)})$ abgesehen von der Reihenfolge mit $F_1, \dots, F_n$ zusammenfallen]. Danach sind die Werte $F_1, \dots, F_n$ vollkommen bestimmt als das kleinste mögliche System von n Werten der Funktion $f(x_1, \dots, x_n)$ für n unabhängige Gitterpunkte.

In dem fünften Kapitel meines Buches „Geometrie der Zahlen“ (I. Heft, Leipzig, 1896) nun habe ich nachgewiesen, daß stets die Ungleichung

$$F_1 \dots F_n J \leq 2^n$$

gilt, wo J das Volumen des Bereichs $f(x_1, \dots, x_n) \leq 1$ ist. Der Beweis dieser Ungleichung erfordert mancherlei Ausführungen. Für die im folgenden beabsichtigten Folgerungen kommt es jedoch nur darauf an, daß sich für $F_1 \dots F_n J$ überhaupt irgendeine, nur von n und nicht weiter von der Funktion $f(x_1, \dots, x_n)$ abhängende obere Grenze angeben läßt. Durch wesentlich einfachere Überlegungen als a. a. O. läßt sich nun folgendes nachweisen.

Hilfssatz I. *Es gilt für den nirgends konkaven Bereich $f(x_1, \dots, x_n) \leq 1$ die Ungleichung:*

$$(12) \quad F_1 \dots F_n J \leq n! 2^n.$$

4. Um den Beweis dieser Ungleichung anzubahnen, ermitteln wir zunächst zu den einmal gewählten n Gitterpunkten $p_1^{(h)}, \dots, p_n^{(h)}$ für $h = 1, \dots, n$ die Substitution (7) und führen damit gewisse neue Variablen $y_1, \dots, y_n$ ein. Es sei K der Körper $f(x_1, \dots, x_n) \leq 1$, und wir wollen für $j = 1, \dots, n$ unter K_j^+ bzw. K_j^- das Gebiet aus K verstehen, für das $y_j \geq 0$, bzw. $y_j \leq 0$ und zudem, (wenn $j < n$ ist), $y_{j+1} = 0, \dots, y_n = 0$ ist; die Vereinigung von K_j^+ und K_j^- heiße K_j ; der Bereich K_n wird nichts anderes als K selbst.

Es sei $y_1 = \delta_1^{(1)}, y_2 = 0, \dots, y_n = 0$ der Punkt (das System $y_1, \dots, y_n$) aus K_1^+ , für den y_1 am größten ist, sodann $y_1 = \delta_1^{(2)}, y_2 = \delta_2^{(2)}, y_3 = 0, \dots, y_n = 0$ ein solcher Punkt aus K_2^+ , für den y_2 möglichst groß ist, usf., schließlich $y_1 = \delta_1^{(n)}, y_2 = \delta_2^{(n)}, \dots, y_n = \delta_n^{(n)}$ ein solcher Punkt aus K_n^+ , für den y_n möglichst groß ist. Dabei sind natürlich $\delta_1^{(1)}, \delta_2^{(2)}, \dots, \delta_n^{(n)}$ positiv. Diese Systeme mögen auch kurz die Punkte $\delta_1, \delta_2, \dots, \delta_n$ und die ihnen entgegengesetzten Systeme die Punkte $-\delta_1, -\delta_2, \dots, -\delta_n$ heißen.

Wir führen nun die Substitution ein:

$$(13) \quad y_1 = \delta_1^{(1)} v_1 + \dots + \delta_1^{(n)} v_n, \dots, y_n = \delta_n^{(n)} v_n \quad (\delta_h^{(k)} = 0, h > k),$$

und wir setzen ihre Determinante $\delta_1^{(1)} \delta_2^{(2)} \dots \delta_n^{(n)} = \Delta$. Für den Punkt δ_j

hat man dann $v_j = 1$, $v_k = 0$ ($k \neq j$); als ein Körper, der den Nullpunkt zum Mittelpunkt hat, enthält K mit $\mathfrak{d}_j$ jedesmal auch den Punkt $-\mathfrak{d}_j$, d. i. $v_j = -1$, $v_k = 0$ ($k \neq j$), und mit diesen $2n$ Systemen $\pm \mathfrak{d}_1, \dots, \pm \mathfrak{d}_n$ enthält K als ein nirgends konkaver Körper (wegen (5)) sogleich den ganzen durch

$$(14) \quad |v_1| + |v_2| + \dots + |v_n| \leq 1$$

definierten Bereich. (Für $n = 3$ stellt dieser Bereich ein Oktaeder vor.)

Es läßt sich nun ein zweiter einfacher Bereich (ein Parallelepipedum) angeben, welcher seinerseits ganz den Körper K in sich enthält. Es habe j einen der Werte $1, \dots, n-1$. Wir suchen einen Ausdruck $\varphi = v_j + \varepsilon^{(j+1)} v_{j+1} + \dots + \varepsilon^{(n)} v_n$ mit geeigneten Konstanten $\varepsilon^{(j+1)}, \dots, \varepsilon^{(n)}$ herzustellen, so daß in K durchweg $\varphi \leq 1$ ist.

Zunächst gilt in K_j : $v_j \leq 1$. Wir bilden nun $\varphi = v_j + \varepsilon v_{j+1}$ mit irgendeiner Konstante ε . In K_j ist $v_{j+1} = 0$, $v_j \leq 1$, also $\varphi \leq 1$. Sowie $\varepsilon \leq -1$ ist, hat man für den Punkt $-\mathfrak{d}_{j+1}$, für den $v_{j+1} = -1$, $v_j = 0$ ist, $\varphi \geq 1$; dann muß in K_{j+1}^+ durchweg $\varphi \leq 1$ sein. Denn hätte man $\varphi > 1$ für irgendeinen Punkt in K_{j+1}^+ , so würde die Strecke von diesem Punkte nach $-\mathfrak{d}_{j+1}$ das Gebiet K_j in einem Punkte treffen, für den ebenfalls $\varphi > 1$ wäre. Sowie andererseits $\varepsilon > 1$ ist, hat man in K_{j+1}^+ für den Punkt $\mathfrak{d}_{j+1}$ jedenfalls $\varphi > 1$. Danach ist das Maximum des Ausdrucks $\varphi = v_j + \varepsilon v_{j+1}$ für ein gegebenes ε im Bereiche K_{j+1}^+ sicher > 1 , wenn $\varepsilon > 1$ ist, und sicher ≤ 1 , wenn $\varepsilon \leq -1$ ist. Dieses Maximum aber ist offenbar eine Funktion von ε , die sich mit ε stetig ändert, und wird es daher einen bestimmten größten Wert $\varepsilon = \varepsilon_j^{(j+1)}$ geben, im Intervalle $-1 \leq \varepsilon \leq 1$ gelegen, für den dieses Maximum noch ≤ 1 , d. h. für den in K_{j+1}^+ noch durchweg $\varphi \leq 1$ ist. Dann ist für jeden Wert ε , der $> \varepsilon_j^{(j+1)}$ ist, in K_{j+1}^+ notwendig irgendwo auch $\varphi > 1$ und daher, weil in K_j überall $\varphi \leq 1$ sein muß, in K_{j+1}^- notwendig überall $\varphi \leq 1$; letzteres muß dann auch noch für den Grenzwert $\varepsilon = \varepsilon_j^{(j+1)}$ gelten, und also ist für diesen Wert im ganzen Bereich K_{j+1} stets $\varphi \leq 1$.

Falls auch noch $j+1 < n$ ist, betrachten wir den neuen Ausdruck $\varphi = v_j + \varepsilon_j^{(j+1)} v_{j+1} + \varepsilon v_{j+2}$ mit irgendeiner Konstante ε . In K_{j+1} ist $v_{j+2} = 0$ und also stets $\varphi \leq 1$. Für ein $\varepsilon \leq -1$ und den Punkt $-\mathfrak{d}_{j+2}$ in K_{j+2}^- ist $\varphi \geq 1$ und daher in K_{j+2}^+ notwendig überall $\varphi \leq 1$. Dagegen ist für ein $\varepsilon > 1$ in K_{j+2}^+ insbesondere $\varphi > 1$ für den Punkt $\mathfrak{d}_{j+2}$. Nunmehr wird es einen bestimmten größten Wert $\varepsilon = \varepsilon_j^{(j+2)}$ geben, im Intervalle $-1 \leq \varepsilon \leq 1$ gelegen, für den in K_{j+2}^+ noch überall $\varphi \leq 1$ ist. Dann ist für jeden Wert $\varepsilon > \varepsilon_j^{(j+2)}$ in K_{j+2}^+ irgendwo $\varphi > 1$ und daher in K_{j+2}^- überall $\varphi \leq 1$, und dieses letztere muß schließlich auch für den Grenzwert $\varepsilon = \varepsilon_j^{(j+2)}$ gelten. Mithin ist $\varphi = v_j + \varepsilon_j^{(j+1)} v_{j+1} + \varepsilon_j^{(j+2)} v_{j+2}$ in K_{j+2} durchweg ≤ 1 .

Es ist nun klar, wie man fortzuschreiten hat, wenn noch $j + 2 < n$ ist, und daß man schließlich zu einem Ausdrucke

$$\varphi_j = v_j + \varepsilon_j^{(j+1)} v_{j+1} + \cdots + \varepsilon_j^{(n)} v_n$$

mit bestimmten Koeffizienten $\varepsilon_j^{(j+1)}, \dots, \varepsilon_j^{(n)}$ kommen wird von solcher Art, daß in K_n , d. i. in K überall $\varphi_j \leq 1$ ist. Weil K ein Körper mit dem Nullpunkt als Mittelpunkt ist, nimmt $-\varphi_j$ in K dieselbe Wertmenge an wie φ_j , und also ist dann in K auch $-\varphi_j \leq 1$, d. h. $\varphi_j \geq -1$.

Endlich hat man noch für $\varphi_n = v_n$ in K stets $\pm \varphi_n \leq 1$.

Man kann auf solche Weise n Formen herstellen:

$$(15) \quad \varphi_1 = v_1 + \varepsilon_1^{(2)} v_2 + \cdots + \varepsilon_1^{(n)} v_n, \quad \varphi_2 = v_2 + \cdots + \varepsilon_2^{(n)} v_n, \dots, \varphi_n = v_n,$$

so daß K ganz in dem Bereiche

$$(16) \quad -1 \leq \varphi_1 \leq 1, \quad -1 \leq \varphi_2 \leq 1, \dots, -1 \leq \varphi_n \leq 1$$

enthalten ist.

Nun ist über diesen Bereich ausgedehnt das Integral $\int dx_1 \dots dx_n = \int dy_1 \dots dy_n = \Delta \cdot \int dv_1 \dots dv_n = \Delta \cdot \int d\varphi_1 \dots d\varphi_n = 2^n \Delta$, wo Δ die Determinante der Gleichungen (13) bedeutet, und also folgt, weil K in diesem Bereiche enthalten ist:

$$(17) \quad J \leq 2^n \Delta.$$

Andererseits ist für den Bereich (14) das Volumen $\int dx_1 \dots dx_n = \Delta \cdot \int dv_1 \dots dv_n = \frac{2^n}{n!} \Delta$ und hat man, weil dieser Bereich (14) ganz in K enthalten ist,

$$(18) \quad \frac{2^n}{n!} \Delta \leq J.$$

5. Weil der Bereich (14) in K enthalten ist, hat man auf der Begrenzung von K , d. h. wenn $f(x_1, \dots, x_n) = 1$ ist, $|v_1| + \cdots + |v_n| \geq 1 = f(x_1, \dots, x_n)$. Wegen (2) und der entsprechenden Regel für den Ausdruck $|v_1| + \cdots + |v_n|$ gilt dann allgemein für jedes beliebige System $x_1, \dots, x_n$:

$$(19) \quad |v_1| + \cdots + |v_n| \geq f(x_1, \dots, x_n).$$

Wir bilden für ein beliebiges System $x_1, \dots, x_n$ unter Benutzung der Substitutionen (7) und (13) den Ausdruck

$$(20) \quad \left| \frac{v_1}{F_1} \right| + \cdots + \left| \frac{v_n}{F_n} \right| = \psi(x_1, \dots, x_n).$$

Ist $x_1, \dots, x_n$ ein vom Nullpunkte verschiedener Gitterpunkt und für ihn unter seinen Zahlen $y_1, \dots, y_n$ etwa y_j die letzte von Null verschiedene, so ist nach der Bedeutung der Werte $F_1, \dots, F_n$ für ihn $f(x_1, \dots, x_n) \geq F_j$; alsdann folgt aus (19), indem, wenn $j < n$ ist, für ihn $v_{j+1}, \dots, v_n$ Null sind, $|v_1| + \cdots + |v_j| \geq F_j$ und mit Rücksicht auf (11) weiter

$$\psi(x_1, \dots, x_n) = \left| \frac{v_1}{F_1} \right| + \cdots + \left| \frac{v_j}{F_j} \right| \geq \frac{|v_1| + \cdots + |v_j|}{F_j} > 1.$$

Danach kann sich für keinen einzigen Gitterpunkt außer dem Nullpunkte $\psi(x_1, \dots, x_n) < 1$ herausstellen. Der durch

$$(21) \quad \psi(x_1, \dots, x_n) \leq 1$$

definierte nirgends konkave Bereich (für $n = 3$ ein Oktaeder) enthält also außer dem Nullpunkte keinen Gitterpunkt im Inneren (d. h. die Begrenzung ausgenommen). Das Integral $\int dx_1 \dots dx_n$ über diesen Bereich ist $= F_1 \dots F_n \cdot \frac{2^n \Delta}{n!}$.

Sind $x_1 = 2r_1, \dots, x_n = 2r_n$ und $x_1 = 2s_1, \dots, x_n = 2s_n$ irgend zwei verschiedene Systeme von je n geraden ganzen Zahlen, so können daher die zwei ihnen entsprechenden Bereiche

$$\psi(x_1 - 2r_1, \dots, x_n - 2r_n) \leq 1, \quad \psi(x_1 - 2s_1, \dots, x_n - 2s_n) \leq 1$$

niemals einen inneren Punkt gemein haben. Denn da man analog zu (4) und (1)

$$2\psi(s_1 - r_1, \dots, s_n - r_n) \leq \psi(x_1 - 2r_1, \dots, x_n - 2r_n) + \psi(2s_1 - x_1, \dots, 2s_n - x_n) \\ \psi(2s_1 - x_1, \dots, 2s_n - x_n) = \psi(x_1 - 2s_1, \dots, x_n - 2s_n)$$

hat, würde in solchem Falle $2\psi(s_1 - r_1, \dots, s_n - r_n) < 2$ hervorgehen, läge also der vom Nullpunkte verschiedene Gitterpunkt $s_1 - r_1, \dots, s_n - r_n$ im Inneren des Bereiches (21).

Man hat im Bereiche (21) stets $|v_j| \leq F_j \leq G$, sodann, wenn δ den größten Betrag unter den Beträgen der $\delta_h^{(j)}$ in (13) ist, $|y_h| \leq n\delta G$, und wenn a den größten Betrag unter den Beträgen der $a_k^{(h)}$ in (7) ist, weiter $|x_k| \leq n^2 a \delta G$, welche Grenze d heiße. In $\psi(x_1 - 2r_1, \dots, x_n - 2r_n) \leq 1$ gilt dann, wenn r der größte Betrag unter denen von $r_1, \dots, r_n$ ist, $|x_k - 2r_k| \leq d$, $|x_k| \leq 2r + d$.

Nun sei Ω irgendeine positive ganze Zahl und man betrachte alle diejenigen Gitterpunkte mit geraden Koordinaten $2r_1, \dots, 2r_n$, für welche jede der Größen r_k einen der Werte $0, \pm 1, \dots, \pm \Omega$ hat. Von den zugehörigen Bereichen $\psi(x_1 - 2r_1, \dots, x_n - 2r_n) \leq 1$ haben keine zwei einen inneren Punkt gemein. Sie liegen alle im Würfel

$$-(2\Omega + d) \leq x_k \leq (2\Omega + d) \quad (k = 1, \dots, n)$$

eingeschlossen. Die Anzahl dieser Bereiche ist $(2\Omega + 1)^n$ und jeder hat dasselbe Volumen wie der Bereich (21). Danach folgt

$$(22) \quad (2\Omega + 1)^n \cdot F_1 \dots F_n \frac{2^n}{n!} \Delta \leq (4\Omega + 2d)^n.$$

Läßt man hierin die ganze Zahl Ω unbegrenzt wachsen, so geht

$$(23) \quad F_1 \dots F_n \frac{2^n}{n!} \Delta \leq 2^n$$

hervor, d. h. das Volumen des Bereichs (21) muß $\leq 2^n$ sein. Vermittels (17) folgt daraus in der Tat der zu erweisende Hilfssatz I.

6. Ferner ist von Interesse:

Hilfssatz II. Die Determinante $|p_k^{(h)}|$ für n unabhängige Gitterpunkte $p_1^{(h)}, \dots, p_n^{(h)}$ ($h=1, \dots, n$), welche $f(p_1^{(h)}, \dots, p_n^{(h)}) = F_h$ ergeben, ist stets dem Betrage nach $\leq n!$.

Nämlich auch der durch

$$(24) \quad |z_1| + \dots + |z_n| \leq 1$$

definierte Bereich enthält keinen Gitterpunkt außer dem Nullpunkte im Inneren. Denn ist $x_1, \dots, x_n$ ein vom Nullpunkte verschiedener Punkt im Inneren dieses Bereichs und von den Größen $z_1, \dots, z_n$ für ihn z_j die letzte von Null verschiedene, so folgt aus (6) vermöge (4), (1), (2)

$$\begin{aligned} f(x_1, \dots, x_n) &\leq |z_1| f(p_1^{(1)}, \dots, p_n^{(1)}) + \dots + |z_j| f(p_1^{(j)}, \dots, p_n^{(j)}) \\ &\leq (|z_1| + \dots + |z_j|) F_j < F_j; \end{aligned}$$

nun ist auch unter den Größen $y_1, \dots, y_n$ für ihn y_j die letzte von Null verschiedene und kann hiernach der Punkt kein Gitterpunkt sein. Auf Grund derselben Überlegungen wie vorhin für den analogen Bereich (21) folgt nunmehr, daß auch das Volumen von (24) sicher $\leq 2^n$ ist. Dieses Volumen ist gleich $\frac{2^n}{n!}$, multipliziert in den Betrag der Determinante $|p_k^{(h)}|$; mithin ist dieser Betrag $\leq n!$.

Das gleiche Resultat ergibt sich bereits in folgender Weise. Der Betrag der Determinante $|p_k^{(h)}|$ ist nach (6), (7), (8) gleich dem reziproken Wert des Produktes $\gamma_1^{(1)} \dots \gamma_n^{(n)}$. Für den Gitterpunkt $p_1^{(h)}, \dots, p_n^{(h)}$ ist $z_h = 1$, $z_k = 0$ ($k \neq h$), also $y_h = \frac{1}{\gamma_h^{(h)}}$, $y_k = 0$ ($k > h$). Da nun dieser Punkt zum Bereiche $f\left(\frac{x_1}{F_h}, \dots, \frac{x_n}{F_h}\right) \leq 1$ gehört, muß für ihn nach der Bedeutung der Größen $\delta_h^{(h)}$ in 4. sich $\frac{y_h}{F_h} \leq \delta_h^{(h)}$ erweisen. Also ist $\frac{1}{\gamma_h^{(h)}} \leq F_h \delta_h^{(h)}$, mithin $\frac{1}{\gamma_1^{(1)} \dots \gamma_n^{(n)}} \leq F_1 \dots F_n \Delta$, woraus vermöge (23) der Hilfssatz II hervorgeht.

Nach (6), (7), (8) sind die Koeffizienten $\gamma_h^{(k)}$ in (8) rationale Zahlen mit der Determinante $|p_k^{(h)}|$ als Nenner, also mit einem Generalnenner $\leq n!$. Nach den Ungleichungen (9) kommen daher für diese Koeffizienten, also für die ganze Substitution $P^{-1}A$ von vornherein eine nur von n abhängende Anzahl von möglichen Ausdrücken in Betracht.

7. Endlich fügen wir die folgende Bemerkung hinzu. Für das System $x_k = a_k^{(h)}$ ($k=1, \dots, n$) in (7) ist $x_k = \gamma_1^{(h)} p_k^{(1)} + \dots + \gamma_h^{(h)} p_k^{(h)}$. Indem nun die Größen $\gamma_1^{(h)}, \dots, \gamma_h^{(h)}$ sämtlich ≥ 0 und ≤ 1 sind, folgt

aus (4) und (2):

$$f(a_1^{(h)}, \dots, a_n^{(h)}) \leq f(p_1^{(1)}, \dots, p_n^{(1)}) + \dots + f(p_1^{(h)}, \dots, p_n^{(h)}) \leq h F_h.$$

Berücksichtigt man nun (12), so zeigt sich, daß für die n ganzzahligen Systeme $a_1^{(h)}, \dots, a_n^{(h)}$ ($h = 1, \dots, n$), deren Determinante $|a_k^{(h)}| = \pm 1$ ist, die Ungleichung

$$(25) \quad f(a_1^{(1)}, \dots, a_n^{(1)}) \dots f(a_1^{(n)}, \dots, a_n^{(n)}) \leq (n!)^2 2^n \frac{1}{J}$$

erfüllt ist.

§ 2. Ein Kriterium für die algebraischen Zahlen n^{ten} Grades.

8. Eine reelle oder komplexe Größe a heißt eine *algebraische Zahl* und zwar n^{ten} Grades, wenn sie einer algebraischen Gleichung n^{ten} Grades

$$x_1 + x_2 a + \dots + x_n a^{n-1} + x_{n+1} a^n = 0 \quad (x_{n+1} \neq 0)$$

mit rationalen ganzzahligen Koeffizienten $x_1, x_2, \dots, x_n, x_{n+1}$, deren letzter $\neq 0$ ist, und nicht bereits einer Gleichung derselben Art von niedrigerem Grade genügt. Dieses ist die *Definition* der algebraischen Zahlen n^{ten} Grades. Im folgenden will ich nun ein *Theorem* entwickeln, wonach die Entscheidung, ob eine gegebene reelle oder komplexe Größe a eine algebraische Zahl n^{ten} Grades ist oder nicht, bereits durch Betrachtung der Werte des Ausdrucks

$$(26) \quad \xi = x_1 + x_2 a + \dots + x_n a^{n-1}$$

für rationale ganze Zahlen $x_1, x_2, \dots, x_n$ geliefert werden kann.

Es sei r irgendeine positive ganze Zahl, und wir lassen für $x_1, \dots, x_n$ in (26) alle diejenigen Systeme von ganzen Zahlen zu, wobei jede Zahl dem Betrage nach $\leq r$ ist, also der Reihe $0, \pm 1, \pm 2, \dots, \pm r$ angehört, mit Ausschluß des einen Systems $x_1 = 0, \dots, x_n = 0$. Unter den betreffenden Systemen kommen gewiß Vereine von n unabhängigen (d. h. mit nichtverschwindender Determinante $|x_k^{(h)}|$) vor; z. B. sind die n Systeme $x_h^{(h)} = 1, x_k^{(h)} = 0$ ($k \neq h$) für $h = 1, \dots, n$ unabhängig. Wir suchen unter allen jenen Systemen ein erstes $x_1 = p_1^{(1)}, \dots, x_n = p_n^{(1)}$ aus, so daß der Betrag von ξ möglichst klein ausfällt. Da wir dabei jedenfalls die Wahl zwischen Paaren entgegengesetzter Systeme haben, wollen wir das System noch derart voraussetzen, daß die letzte von Null verschiedene der Zahlen $p_k^{(1)}$ größer als Null ist. Es sei für das System $\xi = \alpha_1$. Sodann wählen wir unter jenen Systemen ein zweites, von $p_1^{(1)}, \dots, p_n^{(1)}$ unabhängiges System $x_1 = p_1^{(2)}, \dots, x_n = p_n^{(2)}$ aus, wofür nächst dem der Betrag von ξ möglichst klein ausfällt, und es sei wieder unter den Zahlen $p_k^{(2)}$ die letzte nichtverschwindende > 0 . Für dieses zweite System sei $\xi = \alpha_2$. Wir fahren so fort, bis wir schließlich unter jenen Systemen ein n^{tes} von den früheren unabhängiges System $x_1 = p_1^{(n)}, \dots,$

$x_n = p_n^{(n)}$ erlangen, wofür wieder der Betrag von ξ möglichst klein ist, und es sei auch von den Zahlen $p_k^{(n)}$ die letzte von Null verschiedene > 0 . Für dieses n^{te} System sei $\xi = \alpha_n$. Dann hat endlich die Substitution P :

$$(27) \quad x_k = p_k^{(1)} z_1 + p_k^{(2)} z_2 + \cdots + p_k^{(n)} z_n \quad (k = 1, 2, \dots, n)$$

eine von Null verschiedene Determinante, und es geht ξ durch P in

$$(28) \quad \xi P = \chi = \alpha_1 z_1 + \alpha_2 z_2 + \cdots + \alpha_n z_n$$

über. Dabei ist jedenfalls

$$(29) \quad |\alpha_1| \leq |\alpha_2| \leq \cdots \leq |\alpha_n|.$$

Unter speziellen Umständen in bezug auf die Größe a ist es möglich, daß die Systeme $p_1^{(h)}, \dots, p_n^{(h)}$ durch die angegebenen Bedingungen noch nicht eindeutig bestimmt sind, die Festsetzung von P für das gegebene r also noch in mehrfacher Weise geschehen kann. Durch entsprechende Betrachtungen wie in 3. aber erkennt man, daß jedenfalls das System der n absoluten Beträge $|\alpha_1|, |\alpha_2|, \dots, |\alpha_n|$ durch die Zahl r völlig eindeutig bestimmt ist. Es heie P eine zur Zahl r gehörende Substitution.

Man bilde nun eine zur Zahl $r_1 = 1$ gehörende Substitution P_1 . Diese kann auch noch zu $r = 2, 3, \dots$ gehören. Gehört sie nicht zu jeder Zahl r , so sei der größte Wert r , zu dem sie gehört, $r = r_2 - 1$ (wo $r_2 \geq 2$ ist). Es sei sodann P_2 eine zu r_2 gehörende Substitution, und sie gehöre zu den Zahlen r , die $\geq r_2$ und $< r_3$ sind. Sodann sei P_3 eine zu r_3 gehörende Substitution usf. Wir wollen noch die Festsetzung treffen, daß, wenn für eine dieser Substitutionen P_i in der zugehörigen Form $\xi P_i = \chi$ ein Teil der Koeffizienten $\alpha_1, \dots, \alpha_n$ (etwa $\alpha_1, \dots, \alpha_j = 0$ wird, die betreffenden Vertikalreihen $p_1^{(h)}, \dots, p_n^{(h)}$ ($h = 1, \dots, j$) für alle folgenden Substitutionen P_x ($x > i$) unverändert beibehalten werden sollen. Die so entstehende (sei es abbrechende, sei es unendliche) Reihe von Substitutionen $P_1, P_2, P_3, \dots$ soll die zu a gehörende Kette von Substitutionen heißen.

Es sei allgemein $\chi_i = \alpha_1^{(i)} z_1 + \cdots + \alpha_n^{(i)} z_n$ die Form, in welche ξ durch P_i übergeht. Man hat für zwei aufeinanderfolgende Substitutionen P_i, P_{i+1} der Kette:

$$(30) \quad |\alpha_1^{(i+1)}| \leq |\alpha_1^{(i)}|, \dots, |\alpha_n^{(i+1)}| \leq |\alpha_n^{(i)}| \quad (i = 1, 2, \dots),$$

wo jedenfalls nicht alle n Gleichheitszeichen auf einmal gelten können; denn sonst würde ja P_i auch zur Zahl r_{i+1} gehören, während sie nur zu den Werten $r \geq r_i$ und $\leq r_{i+1} - 1$ gehört. In P_i sind jedesmal die Beträge aller Koeffizienten $\leq r_i$ und ist wenigstens einer darunter $> r_i - 1$, also eben $= r_i$. Man erkennt nach diesen Umständen, daß die Reihe der Zahlen $r_1, r_2, r_3, \dots$ eine durch die Größe a völlig bestimmte ist.

9. Ohne an die Aufgabe der einfachsten sukzessiven Ermittlung der Kettenglieder näher heranzutreten, beweise ich nur den folgenden Satz, der bei der Behandlung dieser Aufgabe die wesentlichsten Dienste leistet.

Für eine jede Substitution der Kette ist die Determinante dem Betrage nach $\leq n!$.

Es sei P in (27) eine Substitution der Kette, zu einer Zahl r gehörend, und χ in (28) die Form, in welche ξ durch P übergeht. Dann kann der durch

$$(31) \quad |z_1| + \dots + |z_n| \leq 1$$

definierte Bereich im Inneren (d. h. soweit das Zeichen $<$ gilt), keinen Gitterpunkt außer dem Nullpunkte enthalten. Denn ist $z_1, \dots, z_n$ irgendein, von $0, \dots, 0$ verschiedenes System im Inneren dieses Bereichs (31) und von diesen Größen z_j die letzte von Null verschiedene, so hat man dafür

$$\xi = \alpha_1 z_1 + \dots + \alpha_j z_j, \quad x_k = p_k^{(1)} z_1 + \dots + p_k^{(j)} z_j.$$

Ist *erstens* $|\alpha_j| > 0$, so folgt $|\xi| \leq |\alpha_j|(|z_1| + \dots + |z_j|) < |\alpha_j|$, $|x_k| \leq r(|z_1| + \dots + |z_j|) < r$; alsdann ist das betreffende System $x_1, \dots, x_n$ kein Gitterpunkt, denn für einen Gitterpunkt, bei dem die Beträge $|x_k| \leq r$ sind und $z_j \neq 0$ ist, muß $|\xi| \geq |\alpha_j|$ sein.

Ist *zweitens* $\alpha_j = 0$, so sei P_x die erste Substitution der Kette, für welche sich $\alpha_1^{(x)} = \dots = \alpha_j^{(x)} = 0$ findet, während, wenn $x > 1$ ist, in der zu P_{x-1} gehörenden Form χ_{x-1} der erste nichtverschwindende Koeffizient $\alpha_{j'}^{(x-1)}$ mit einem Index $j' \leq j$ sei. Dann ist also r_x die kleinste Zahl, für welche die ersten j' Vertikalreihen einer zugehörigen Substitution sämtlich $\xi = 0$ machen. Nun gehören zufolge einer in 8. getroffenen Festsetzung die ersten j Vertikalreihen von P bereits zu P_x und, wenn $x > 1$ und $j' > 1$ ist, die ersten $j' - 1$ Vertikalreihen von P bereits zu P_{x-1} . Für das System $z_1, \dots, z_n$ folgt jetzt $\xi = 0$, $|x_k| \leq r_x(|z_1| + \dots + |z_j|) < r_x$. Im Falle $x = 1$ ist $r_1 = 1$ und kann also $x_1, \dots, x_n$ hier kein Gitterpunkt sein. Ist aber $x > 1$ und wäre $x_1, \dots, x_n$ hier ein Gitterpunkt, so wäre derselbe, da $z_j \neq 0$ ist, ja vom Nullpunkte verschieden und zudem, falls $j' > 1$ ist, auch von den Gitterpunkten in den ersten $j' - 1$ Vertikalreihen von P_{x-1} unabhängig; also würde bereits zu einer gewissen Zahl $< r_x$ eine Substitution gehören müssen, in welcher mindestens j' Vertikalreihen sämtlich $\xi = 0$ machen.

Aus dem Umstande, daß (31) keinen Gitterpunkt im Inneren enthält, folgt wie im ersten Beweise des Hilfssatzes II (s. 6.), daß die Determinante $|p_k^{(n)}|$ von P dem Betrage nach $\leq n!$ ist.

10. Es sei $n > 1$. Ferner wollen wir von dem Falle absehen, daß $n = 2$ und a komplex ist; alsdann würde $|\xi| = |x_1 + ax_2|$ unter einer

gegebenen Grenze nur für eine endliche Anzahl von ganzzahligen Systemen x_1, x_2 liegen, wäre also die Substitutionenkette zu a jedenfalls eine abbrechende; andererseits ist eine komplexe Größe $a = b + ic$ dann und nur dann eine algebraische Zahl zweiten Grades, wenn b sowie c^2 rational sind.

Auf Grund der in 8. entwickelten Begriffe entsteht nun folgendes vollständige

Kriterium für die algebraischen Zahlen n^{ten} Grades:

Es sei a eine beliebige reelle oder komplexe Größe, im ersten Falle $\sigma = 1$, im zweiten $\sigma = 2$ und $n > \sigma$. Es sei

$$\xi = x_1 + ax_2 + \dots + a^{n-1}x_n,$$

sodann $P_1, P_2, P_3, \dots$ die zu a gehörige Kette von Substitutionen mit n Variablen, und $\chi_1, \chi_2, \chi_3, \dots$ seien die Formen, in welche ξ durch $P_1, P_2, P_3, \dots$ übergeht.

1° Ist a nicht eine algebraische Zahl n^{ten} oder niederen Grades, so bricht die Kette niemals ab, und alle Gleichungen $\chi_1 = 0, \chi_2 = 0, \dots$ sind verschieden (d. h. keine zwei der Formen $\chi_1, \chi_2, \dots$ unterscheiden sich bloß durch einen Faktor). In jeder Form χ_x sind alle Koeffizienten von Null verschieden.

2° Ist a eine algebraische Zahl n^{ten} Grades, so bricht die Kette niemals ab, unter den Gleichungen $\chi_1 = 0, \chi_2 = 0, \dots$ kommen nur eine **endliche** Anzahl verschiedener vor (d. h. alle diese unendlich vielen Formen entstehen aus einer endlichen Anzahl unter ihnen durch Multiplikation mit Faktoren), in jeder Form χ_x sind alle Koeffizienten von Null verschieden.

3° Ist a eine algebraische Zahl $n - m^{\text{ten}}$ Grades, wo $m > 0$ und $n - m > \sigma$ ist, so bricht die Kette niemals ab, unter den Gleichungen $\chi_1 = 0, \chi_2 = 0, \dots$ kommen nur eine endliche Anzahl verschiedener vor, in den Formen χ_x sind von einer gewissen an stets die m ersten Koeffizienten $= 0$, die übrigen $n - m$ sind beständig von Null verschieden.

4° Ist a eine algebraische Zahl σ^{ten} Grades (also reell und rational oder komplex und vom zweiten Grade), so bricht die Kette nach einer endlichen Anzahl von Gliedern ab.

11. Wir beginnen den Nachweis dieses Satzes mit der Feststellung, daß, wenn zwei der Gleichungen $\chi_1 = 0, \chi_2 = 0, \dots$ identisch sind, die Größe a notwendig eine algebraische Zahl n^{ten} oder niederen Grades ist.

Es seien also χ_i und χ_x zwei unter jenen Formen, die sich bloß durch einen Faktor θ unterscheiden, so daß $\chi_x = \theta\chi_i$ ist. Es sei $i < x$, so ist zufolge der Bemerkungen bei (30) der Betrag $|\theta| < 1$. Es seien P_i und P_x die Substitutionen der Kette, welche ξ in χ_i und χ_x überführen. Durch $P_x P_i^{-1}$ geht dann ξ in $\theta\chi_i P_i^{-1}$, d. i. in $\theta\xi$ über. Es sei

$$(32) \quad |q_h^{(k)}| \quad (h, k=1, \dots, n),$$

h als den Index der Horizontal-, k als den Index der Vertikalreihen gedacht, das Koeffizientensystem der Substitution $P_x P_i^{-1}$, so ergibt die Vergleichung der Koeffizienten in den Formen $\theta \xi$ und $\xi P_x P_i^{-1}$:

$$(33) \quad \theta a^{k-1} = q_1^{(k)} + a q_2^{(k)} + \dots + a^{n-1} q_n^{(k)} \quad (k=1, \dots, n).$$

Die Koeffizienten $q_h^{(k)}$ sind sämtlich rationale Zahlen. Man entnimmt aus diesen Gleichungen, daß die Determinante

$$(34) \quad |\theta e_h^{(k)} - q_h^{(k)}| = 0 \quad \left(\begin{array}{l} e_h^{(k)} = 0, \quad h \neq k; \quad e_h^{(h)} = 1 \\ h, k = 1, \dots, n \end{array} \right)$$

ist, eine Gleichung, die symbolisch $|\theta P_i P_i^{-1} - P_x P_i^{-1}| = 0$ geschrieben werden kann.

Nimmt man zwei aufeinanderfolgende der Gleichungen (33), für $k=j$ und $k=j+1$ ($j=1, \dots, n-1$), so entsteht aus ihnen durch Elimination von θ :

$$(35) \quad q_1^{(j+1)} + a(q_2^{(j+1)} - q_1^{(j)}) + \dots + a^{n-1}(q_n^{(j+1)} - q_n^{(j)}) - a^n q_n^{(j)} = 0.$$

Man hat $n-1$ solcher Gleichungen für $j=1, \dots, n-1$.

Entweder ist nun wenigstens eine unter ihnen so beschaffen, daß in ihr nicht alle Koeffizienten der Potenzen von a Null sind; dann erweist sich durch die betreffende Gleichung die Größe a als eine algebraische Zahl n^{ten} oder niederen Grades.

Oder aber, es wären die Ausdrücke in a auf den linken Seiten der Gleichungen (35) für $j=1, \dots, n-1$ sämtlich identisch gleich Null; dann hätte man

$$(36) \quad q_1^{(j+1)} = 0, \quad q_2^{(j+1)} = q_1^{(j)}, \dots, q_n^{(j+1)} = q_n^{(j)}, \quad 0 = q_n^{(j)}$$

für $j=1, \dots, n-1$, also zuvörderst

$$q_1^{(2)} = 0, \quad q_1^{(3)} = 0, \dots, q_1^{(n)} = 0; \quad q_n^{(1)} = 0, \quad q_n^{(2)} = 0, \dots, q_n^{(n-1)} = 0,$$

sodann allgemein, wenn $k > h$ ist, $q_h^{(k)} = q_{h-1}^{(k-1)} = \dots = q_1^{(k-h+1)} = 0$, und wenn $k < h$ ist, $q_h^{(k)} = q_{h+1}^{(k+1)} = \dots = q_n^{(n-h+k)} = 0$, endlich noch $q_1^{(1)} = q_2^{(2)} = \dots = q_n^{(n)}$, welcher Wert $= q$ gesetzt werde. Nunmehr würden die Gleichungen (33) in $\theta a^{k-1} = q a^{k-1}$ ($k=1, \dots, n$) übergehen; man fände $\theta = q$ und hätte allgemein $q_h^{(k)} = q e_h^{(k)}$, also $P_x P_i^{-1} = q P_i P_i^{-1}$, mithin $P_x = \theta P_i$; jeder Koeffizient in P_x wäre gleich dem entsprechenden Koeffizienten aus P_i , multipliziert in θ . Nun hat von den Koeffizienten in P_x wenigstens einer einen Betrag $= r_x$ und die Beträge der Koeffizienten in P_i sind sämtlich $\leq r_i$; es würde somit $r_x \leq |\theta| r_i$ folgen, was mit $r_x > r_i$ und $|\theta| < 1$ im Widerspruch stünde. Die zweite Annahme, daß keine der Gleichungen (34) eine Bedingung für a gibt, ist danach unzulässig, und mithin ist notwendig a eine algebraische Zahl n^{ten} oder niederen Grades.

12. Wir ziehen jetzt die Ergebnisse des § 1 heran. Es sei ε eine beliebige positive GröÙe; wir nehmen für die linearen Formen $\xi_1, \xi_2, \dots$, an welche die Betrachtungen in § 1 anknüpfen, die $n + 1$ Formen

$$x_1, \dots, x_n \text{ und } \frac{\xi}{\varepsilon} = \frac{1}{\varepsilon} (x_1 + ax_2 + \dots + a^{n-1}x_n).$$

Der Körper K wird dann der durch

$$(37) \quad -1 \leq x_1 \leq 1, \dots, -1 \leq x_n \leq 1, \quad |\xi| \leq \varepsilon$$

bestimmte Bereich. In diesem Bereich ist offenbar stets

$$|\xi| \leq 1 + |a| + \dots + |a^{n-1}|,$$

welche GröÙe C^* heiÙe; wir wollen deshalb jedenfalls $\varepsilon \leq C^*$ voraussetzen. Wir haben nach (12) die fundamentale Ungleichung

$$(12b) \quad F_1 \dots F_n J \leq n! 2^n,$$

wo J das Volumen von K ist und die Bedeutung von $F_1, \dots, F_n$ aus 3. zu ersehen ist.

Es werde $\sigma = 1$ gesetzt, wenn a reell ist, $\sigma = 2$, wenn a komplex ist, und im letzteren Falle setze man $\xi = \eta + i\xi$, so daÙ η und ξ Formen mit reellen Koeffizienten sind. Die Bedingung $|\xi| \leq \varepsilon$ kommt für $\sigma = 1$ auf $-\varepsilon \leq \xi \leq \varepsilon$, für $\sigma = 2$ auf $\eta^2 + \xi^2 \leq \varepsilon^2$ hinaus. Denkt man sich zu ξ , bzw. zu η, ξ weitere $n - \sigma$ reelle lineare Formen $v_1, \dots, v_{n-\sigma}$ hinzugenommen, so daÙ n Formen mit einer Determinante 1 entstehen, so erkennt man, daÙ mit nach Null abnehmendem ε das Verhältnis $\frac{J}{2^\sigma \varepsilon}$ für $\sigma = 1$ oder $\frac{J}{\pi \varepsilon^2}$ für $\sigma = 2$ nach dem Werte des über den Bereich

$$|\xi| = 0, \quad -1 \leq x_1 \leq 1, \dots, -1 \leq x_n \leq 1$$

erstreckten Integrals $\int dv_1 \dots dv_{n-\sigma}$, also nach einer bestimmten positiven GröÙe konvergiert. Danach wird man eine endliche von a abhängende, von ε aber *unabhängige* GröÙe M angeben können, so daÙ für alle Werte $\varepsilon \leq C^*$ stets

$$(38) \quad \left(\frac{n! 2^n}{J} \right)^{\frac{1}{\sigma}} \varepsilon \leq M$$

ist. Die Ungleichung (12b) geht dann in

$$(39) \quad \varepsilon (F_1 \dots F_n)^{\frac{1}{\sigma}} \leq M$$

über. Wegen $F_1 \leq \dots \leq F_n$ erhält man daraus insbesondere

$$(40) \quad \varepsilon F_1^{\frac{n}{\sigma}} \leq M$$

und ferner

$$(41) \quad \varepsilon^\sigma F_1^{n-1} F_n \leq M^\sigma.$$

Nunmehr sollen einige Folgerungen aus diesen Ungleichungen (40) und (41) entwickelt werden.

13. Wenn in einer Form χ zur Kette der erste Koeffizient $\alpha_1 = 0$ ausfällt, erweist sich durch diese Gleichung a als eine algebraische Zahl $n - 1^{\text{ten}}$ oder niederen Grades. Eine solche Beschaffenheit von a wollen wir zunächst ausschließen. Dann ist also für jedes Kettenglied gewiß $|\alpha_1| > 0$. Wir zeigen, daß alsdann im Verlauf der Kette der Betrag $|\alpha_1|$ schließlich unter jede Grenze sinken muß.

Es sei P in (27) eine zur Zahl r gehörende Substitution, χ in (28) die zugehörige Transformierte von ξ . Wir nehmen in 12. den Parameter $\varepsilon = \frac{|\alpha_1|}{r}$; (diese Größe ist $\leq C^*$, weil $|p_1^{(1)}|, \dots, |p_n^{(1)}| \leq r$ ist). Man hat hier ein von $0, \dots, 0$ verschiedenes ganzzahliges System $x_1, \dots, x_n$, wofür $|x_k| \leq r$ ($k = 1, \dots, n$) und $\left| \frac{\xi}{\varepsilon} \right| = \left| \frac{\alpha_1}{\frac{|\alpha_1|}{r}} \right| = r$ ist, aber kein System

dieser Art, wofür stets $|x_k| < r$ und $\left| \frac{\xi}{\varepsilon} \right| < r$ wäre. Nach der Bedeutung von F_1 (s. 3.) ist daher dann $F_1 = r$. Die Ungleichung (40) ergibt nunmehr

$$(42) \quad |\alpha_1| \leq M r^{-\frac{n-\sigma}{\sigma}}$$

Diese Abschätzung für $|\alpha_1|$ zeigt in der Tat, daß mit wachsendem r der Wert $|\alpha_1|$ schließlich unter jede Grenze sinken muß.*) In den bezeichneten Fällen, in denen α_1 niemals Null werden kann, hat infolgedessen die Kette notwendig einen unendlichen Verlauf, sie kann niemals abbrechen.

Insbesondere bricht auf diese Weise die Kette niemals ab, wenn a *nicht* eine algebraische Zahl n^{ten} oder niedrigeren Grades ist. Verbindet

*) Sind α, β zwei reelle Größen und ist $0 < |\beta| \leq |\alpha|$, so hat man eine ganze Zahl y , so daß $|\alpha + y\beta| \leq \frac{1}{2}|\beta|$ ist. Sind α, β, γ drei reelle oder komplexe Größen, so daß $0 < |\gamma| \leq |\beta| \leq |\alpha|$ und $\frac{\gamma}{\beta} = \delta + i\varepsilon$ nicht reell, also $\varepsilon \neq 0$ ist, so kann man zwei reelle Größen v, w finden, so daß $\alpha + v\beta + w\gamma = \alpha + (v + \delta w)\beta + iw\varepsilon\beta = 0$ ist, und sind dann z und y zwei ganze Zahlen, so daß $|z - w| \leq \frac{1}{2}$ und $|y + \delta z - v - \delta w| \leq \frac{1}{2}$ ist, so wird $|\alpha + y\beta + z\gamma| \leq \frac{1}{\sqrt{2}}|\beta|$. Mit diesen Hilfsmitteln kann man direkt einsehen, daß in den Formen χ zur Kette sogar $|\alpha_n|$ schließlich unter jede Grenze sinkt, mit Ausnahme des Falles, daß $n = 3$, a komplex und der reelle Teil von a rational ist, in welchem Falle nur $|\alpha_1|$ und $|\alpha_2|$ unter jede Grenze sinken, aber für $|\alpha_3|$ eine positive untere Grenze besteht, und ferner der Fälle, daß a eine algebraische Zahl σ^{ten} Grades ist, in welchen Fällen die Kette mit einem letzten Gliede abbricht.

man hiermit das Ergebnis aus 11., wonach unter derselben Voraussetzung niemals zwei der Formen χ_x sich bloß durch einen Faktor unterscheiden, so ist zunächst der Punkt 1^o unseres Kriteriums völlig sichergestellt.

14. Um den Punkt 2^o zu erweisen, nehmen wir nunmehr an, daß a als eine algebraische Zahl n^{ten} Grades gegeben ist. In diesem Falle können wir außer mit der Ungleichung (39) noch mit dem Satze operieren, daß die Norm einer von Null verschiedenen *ganzen* algebraischen Zahl eine von Null verschiedene *ganze rationale* Zahl und daher dem Betrage nach ≥ 1 ist.

Es sei

$$(43) \quad g_0 a^n + g_1 a^{n-1} + \dots + g_n = 0$$

die Gleichung n^{ten} Grades mit rationalen ganzen Koeffizienten ohne gemeinsamen Teiler und positivem g_0 , der a genügt. Es sei, wenn a komplex, $\sigma = 2$ ist, a^0 die zu a konjugiert imaginäre Größe. Die Wurzeln jener Gleichung n^{ten} Grades außer a , bzw. außer a und a^0 , mögen $a', a'', \dots, a^{(n-\sigma)}$ heißen. Ferner sollen die zu einer Zahl ξ des Körpers von a konjugierten Zahlen in den Körpern von $(a^0), a', \dots, a^{(n-\sigma)}$ mit $(\xi^0), \xi', \dots, \xi^{(n-\sigma)}$ bezeichnet werden.

Das Multiplum $g_0 a$ ist eine *ganze* algebraische Zahl n^{ten} Grades. Infolgedessen ist für ein ganzzahliges, von 0, ..., 0 verschiedenes System $x_1, \dots, x_n$ der Wert von $g_0^{n-1} \xi$ stets eine von Null verschiedene *ganze* Zahl im Körper von a und daher deren Norm in diesem Körper $Nm(g_0^{n-1} \xi) = g_0^{n(n-1)} \xi(\xi^0) \xi' \dots \xi^{(n-\sigma)}$ dem Betrage nach ≥ 1 ; der eingeklammerte Faktor ξ^0 kommt für $\sigma = 2$ in Betracht und dann ist $|\xi^0| = |\xi|$. Es sei C der größte Wert unter den $n - \sigma$ Ausdrücken $1 + |a^{(j)}| + \dots + |a^{(j)}|^{n-1}$ für $j = 1, \dots, n - \sigma$. Sind $x_1, \dots, x_n$ in ihren Beträgen $\leq r$, wo r eine positive Größe ist, so hat man

$$(44) \quad |\xi^{(j)}| \leq Cr \quad (j = 1, \dots, n - \sigma)$$

und entsteht nunmehr die Ungleichung $g_0^{n(n-1)} C^{n-\sigma} |\xi|^{\sigma} r^{n-\sigma} \geq 1$ oder

$$(45) \quad |\xi|^{\sigma} r^{n-\sigma} \geq b^{\sigma},$$

wo b eine gewisse, nur von a , nicht von der Größe von r abhängende Konstante vorstellt.

Es sei wieder P die zu einer ganzen Zahl r gehörende Substitution der Kette und χ in (28) die Form ξP . Dann geht, wenn für $x_1, \dots, x_n$ die erste Vertikalreihe von P genommen wird, aus (45) zuvörderst

$$(46) \quad |\alpha_1| \geq br^{-\frac{n-\sigma}{\sigma}}$$

hervor.

Es sei sodann ε wieder eine beliebige positive Größe $\leq C^*$ und betrachten wir die Ungleichung (41) dafür. Nach der Bedeutung von F_1

in 12. hat man ein von $0, \dots, 0$ verschiedenes ganzzahliges System $x_1, \dots, x_n$, wofür $|x_k| \leq F_1$ ($k=1, \dots, n$) und $\left| \frac{\xi}{\varepsilon} \right| \leq F_1$, also $|\xi| \leq \varepsilon F_1$ ist. Die Ungleichung (45) ergibt daher $(\varepsilon F_1)^\sigma F_1^{n-\sigma} \geq b^\sigma$ oder

$$(47) \quad \varepsilon F_1^{\frac{n}{\sigma}} \geq b.$$

Führt man die hierdurch angewiesene untere Grenze in (41) ein, so folgt

$$(48) \quad \varepsilon F_n^{\frac{n}{\sigma}} \leq B,$$

wo $B = \frac{M^n}{b^{n-1}}$ wieder eine nur von a , nicht von ε abhängende Konstante vorstellt.

Nunmehr wollen wir speziell $\varepsilon = \frac{|\alpha_n|}{r}$ nehmen, wo α_n der letzte Koeffizient in χ ist. Dann ist $F_n = r$; denn man hat hier n unabhängige ganzzahlige Systeme $x_1, \dots, x_n$, für welche $|x_k| \leq r$ ($k=1, \dots, n$) und $\left| \frac{\xi}{\varepsilon} \right| \leq \frac{|\alpha_n|}{\frac{|\alpha_n|}{r}} = r$ ist, aber nach der Bedeutung von $|\alpha_n|$ jedenfalls nicht n unabhängige ganzzahlige Systeme $x_1, \dots, x_n$, für welche $|x_k| < r$ ($k=1, \dots, n$) und $\left| \frac{\xi}{\varepsilon} \right| < r$ wäre. Die Formel (48) ergibt nunmehr

$$(49) \quad |\alpha_n| \leq B r^{-\frac{n-\sigma}{\sigma}}.$$

Wir stellen (46), (49) und (29) zusammen in:

$$(50) \quad b r^{-\frac{n-\sigma}{\sigma}} \leq |\alpha_1| \leq \dots \leq |\alpha_n| \leq B r^{-\frac{n-\sigma}{\sigma}}.$$

Man ersieht daraus zunächst, daß im Verlauf der Kette selbst $|\alpha_n|$ schließlich unter jede Grenze sinkt. Die Kette bricht also jedenfalls niemals ab.

Gewisse weitere Ungleichungen erschließen wir für die Normen der Zahlen α_k ($k=1, \dots, n$) und für ihre konjugierten Zahlen $\alpha_k^{(j)}$. Nach (44) werden wir, da die Zahlen der k^{ten} Vertikalreihe von P , welche $\xi = \alpha_k$ machen, $\leq r$ sind,

$$(51) \quad |\alpha_k^{(j)}| \leq C r \quad (j=1, \dots, n-\sigma)$$

haben. Nehmen wir dazu die in (49) erhaltene obere Grenze für $|\alpha_k|$, womit noch im Falle $\sigma=2$ sich $|\alpha_k^0|$ deckt, so folgt

$$(52) \quad |Nm(g_0^{n-1}\alpha_k)| \leq g_0^{n(n-1)} B^\sigma C^{n-\sigma},$$

wo die Grenze rechts nur von a , nicht von der Zahl r abhängig erscheint. Nun ist $Nm(g_0^{n-1}\alpha_k)$ eine ganze rationale Zahl und daher nach dieser Ungleichung nur einer endlichen Anzahl von Werten bei allen möglichen Werten von r fähig.

Andererseits hat man, indem diese Zahl $Nm(g_0^{n-1}\alpha_k)$ von Null verschieden ist:

$$(53) \quad |Nm(g_0^{n-1}\alpha_k)| = g_0^{n(n-1)} |\alpha_k(\alpha_k^0)\alpha_k' \dots \alpha_k^{(n-\sigma)}| \geq 1.$$

Führt man hierin für $|\alpha_k|$ ($|\alpha_k^0|$) die obere Grenze aus (50) und für die Beträge $|\alpha_k'|, \dots, |\alpha_k^{(n-\sigma)}|$ mit Ausnahme eines Faktors $|\alpha_k^{(j)}|$ die obere Grenze aus (51) ein, so folgt für diesen noch übrigen Betrag:

$$(54) \quad |\alpha_k^{(j)}| \geq cr \quad (j = 1, \dots, n - \sigma),$$

wo $c = \frac{1}{g_0^{n(n-1)} B^\sigma C^{n-\sigma-1}}$ wieder nur von a , nicht von r abhängig ist.

Sind nun h, k irgend zwei der Indizes $1, 2, \dots, n$, so hat man wegen (50):

$$(55) \quad \left| \frac{\alpha_k}{\alpha_h} \right| \left(= \left| \frac{\alpha_k^0}{\alpha_h^0} \right| \right) \leq \frac{B}{b}$$

und aus (51) und (54):

$$(56) \quad \left| \frac{\alpha_k^{(j)}}{\alpha_h^{(j)}} \right| \leq \frac{C}{c} \quad (j = 1, \dots, n - \sigma).$$

Zudem sind $g_0^{n-1}\alpha_h$ und $g_0^{n-1}\alpha_k$ und alle ihre konjugierten Zahlen ganze algebraische Zahlen und liegt nach (52) die ganze rationale Zahl $|Nm(g_0^{n-1}\alpha_h)|$ unter einer bestimmten von a allein abhängenden Grenze. Aus diesen Umständen geht hervor, daß in der algebraischen Gleichung n^{ten} Grades für t :

$$(57) \quad Nm(g_0^{n-1}\alpha_h) \cdot \left(t - \frac{\alpha_k}{\alpha_h}\right) \left(\left(t - \frac{\alpha_k^0}{\alpha_h^0}\right)\right) \left(t - \frac{\alpha_k'}{\alpha_h'}\right) \dots \left(t - \frac{\alpha_k^{(n-\sigma)}}{\alpha_h^{(n-\sigma)}}\right) = 0$$

alle $n + 1$ Koeffizienten der linken Seite ganze rationale Zahlen sind und in ihren Beträgen unterhalb bestimmter nur von a abhängender Grenzen liegen. Danach kommen für die Koeffizienten dieser Gleichung von vornherein für alle r nur eine endliche Anzahl von Wertsystemen und also für ein jedes Verhältnis $\frac{\alpha_k}{\alpha_h}$ von vornherein für alle r nur eine endliche Anzahl von algebraischen Zahlen in Betracht. Mithin können in der Tat die sämtlichen unendlich vielen Linearformen $\chi_1, \chi_2, \dots$ der Kette hier, wo a eine algebraische Zahl n^{ten} Grades ist, nur eine *endliche* Anzahl von verschiedenen Verhältnissen $\alpha_1 : \alpha_2 : \dots : \alpha_n$ der Koeffizienten darbieten, sie gehen also sämtlich aus einer endlichen Anzahl unter ihnen durch Multiplikation mit Faktoren hervor. Damit ist der Punkt 2^o unseres Kriteriums völlig erwiesen.

15. Wir können noch eine Bemerkung über die Natur derjenigen Faktoren θ hinzufügen, die hier in Beziehungen $\chi_x = \theta \chi_i$ auftreten und die als Quotienten von Zahlen im Körper von a jedenfalls auch in diesem Körper liegen.

Zu jeder Substitution P_i der Kette kann man nach der in 2. aus-
einandergesetzten Methode eine ganzzahlige Substitution A_i mit einer
Determinante ± 1 bestimmen, so daß die Koeffizienten in $P_i^{-1}A_i$ den in
(8) und (9) für die Zahlen $\gamma_h^{(k)}$ angegebenen Bedingungen entsprechen.
In dieser Substitution $P_i^{-1}A_i$ sind die Koeffizienten rationale Zahlen mit
der Determinante von P_i als Nenner. Letztere ist nach dem in 9. Be-
wiesenen stets dem Betrage nach $\leq n!$. Nach den Bedingungen (8) und (9)
kommen dadurch für alle Koeffizienten von $P_i^{-1}A_i$ von vornherein nur
eine endliche Anzahl verschiedener Werte in Frage. Verbindet man damit
das soeben in 14. erzielte Resultat, so erkennt man, daß auch für den In-
begriff der Gleichung $\chi_i = 0$ und der Substitution $P_i^{-1}A_i$ zusammen in
der ganzen unendlichen Kette nur eine endliche Anzahl von verschiedenen
Bestimmungen vorkommen.

Hat man nun für zwei Indizes i und x sowohl $\chi_x = \theta \chi_i$, wo θ ein
Faktor ist, als auch $P_x^{-1}A_x = P_i^{-1}A_i$, so wird $P_x P_i^{-1} = A_x A_i^{-1}$, also
eine ganzzahlige Substitution mit einer Determinante ± 1 . Durch diese
Substitution geht die Linearform ξ (in (26)) in $\theta \xi$ über, während $A_i A_i^{-1}$,
d. i. die identische Substitution, ξ in ξ überführt. Für den Faktor θ er-
hält man dadurch eine Gleichung n^{ten} Grades, welche sich in der oben
bei (34) erklärten symbolischen Weise $|\theta A_i A_i^{-1} - A_x A_i^{-1}| = 0$ oder
 $|\theta A_i - A_x| = 0$ schreiben läßt. In dieser Gleichung sind alle Koeffizienten
ganze rationale Zahlen und ist sowohl der erste wie der letzte Koeffizient
gleich ± 1 . Also ist dann sowohl θ wie $\frac{1}{\theta}$ eine ganze algebraische Zahl,
mithin θ eine Einheit in dem Zahlenkörper von a . Sonach gilt der Zusatz:
Unter den Formen $\chi_1, \chi_2, \dots$ der Kette kann man eine endliche Anzahl
angeben so, daß aus ihnen alle Formen der Kette durch Multiplikation
mit solchen Faktoren hervorgehen, welche *Einheiten* im Körper von a sind.

Ist $\chi_i = \alpha_1 z_1 + \dots + \alpha_n z_n$, $\chi_x = \beta_1 z_1 + \dots + \beta_n z_n$ und sind r_i, r_x die
Zahlen, zu denen P_i, P_x gehören, so hat man ferner nach (51) und (54)

$$cr_i \leq |\alpha_k^{(j)}| \leq Cr_i, \quad cr_x \leq |\beta_k^{(j)}| \leq Cr_x \quad \left(\begin{matrix} k=1, \dots, n, \\ j=1, \dots, n-\sigma \end{matrix} \right);$$

indem $\beta_k = \theta \alpha_k$ ist, folgt daraus

$$\frac{c}{C} \frac{r_x}{r_i} \geq |\theta^{(j)}| \leq \frac{C}{c} \frac{r_x}{r_i} \quad (j=1, \dots, n-\sigma).$$

Man hat sodann die von den Zahlen r_i, r_x freie Beziehung

$$(58) \quad \frac{c^2}{C^2} \leq \left| \frac{\theta^{(j)}}{\theta^{(j')}} \right| \leq \frac{C^2}{c^2} \quad (j \neq j'; j, j' = 1, \dots, n-\sigma).$$

Für die Einheiten θ , auf die man hier geführt wird, liegen also die Ver-
hältnisse der Beträge der $n-\sigma$ konjugierten Werte $\theta', \dots, \theta^{(n-\sigma)}$ stets
zwischen zwei von vornherein anzuweisenden Grenzen, ein Umstand, der

für die weitere Erforschung der Substitutionenketten zu algebraischen Zahlen n^{ten} Grades von hervorragender Bedeutung ist.

16. Es sei jetzt a eine algebraische Zahl $n - m^{\text{ten}}$ Grades, wo $m > 0$ und $< n$ ist, und es sei

$$(59) \quad g_{n-m} + g_{n-m-1}a + \dots + g_0 a^{n-m} = 0$$

die Gleichung $n - m^{\text{ten}}$ Grades mit rationalen ganzen Koeffizienten ohne gemeinsamen Teiler und positivem g_0 , der a genügt. Es sei unter den Beträgen der Koeffizienten dieser Gleichung der größte Wert g^* . Man entnimmt aus (59)

$$g_{n-m} a^{h-1} + g_{n-m-1} a^h + \dots + g_0 a^{n-m+h-1} = 0 \quad (h = 1, \dots, m),$$

d. i. $x_1 + ax_2 + \dots + a^{n-1}x_n = 0$ für gewisse m besondere, von $0, \dots, 0$ verschiedene ganzzahlige Systeme $x_1, \dots, x_n$. Diese m Systeme sind voneinander unabhängig, denn schreibt man sie in m Vertikalreihen auf, so hat man in den letzten m Horizontalreihen der entstehenden Matrix, welche die Werte $x_{n-m+1}, \dots, x_n$ enthalten, ein quadratisches Schema, wobei in der Hauptdiagonale alle Elemente $= g_0$ und unterhalb derselben alle Elemente $= 0$ sind, ein Schema also, das die von Null verschiedene Determinante g_0^m ergibt. Alle jene Zahlen $x_1, \dots, x_n$ sind ferner dem Betrage nach $\leq g^*$.

Es sei nun P eine zur Zahl r gehörende Substitution der Kette und χ die Form, in welche ξ durch P übergeht. Aus dem eben Gesagten erkennt man, daß, sowie $r \geq g^*$ ist, in χ jedenfalls $\alpha_1 = 0, \dots, \alpha_m = 0$ sein muß. Dagegen muß stets α_{m+1} von Null verschieden bleiben. Denn hätte man $m + 1$ unabhängige ganzzahlige Systeme $x_1, \dots, x_n$, wofür $\xi = 0$ ausfiele, so würde man aus den betreffenden $m + 1$ Gleichungen durch m Mal hintereinander vorgenommene Elimination der jedesmal sich darbietenden höchsten Potenz von a schließlich eine Gleichung für a mit rationalen Koeffizienten von einem Grade $< n - m$ gewinnen können.

Es sei wieder $\sigma = 1$ oder $= 2$, je nachdem a reell oder komplex ist, und im letzteren Falle sei a^0 die zu a konjugiert imaginäre Zahl. Ferner seien $a', a'', \dots, a^{(n-m-\sigma)}$ die Wurzeln der Gleichung $n - m^{\text{ten}}$ Grades für a außer a , bzw. außer a und a^0 . Endlich, wenn α eine Zahl im Körper von a bedeutet, so seien allgemein $(\alpha^0), \alpha', \dots, \alpha^{(n-m-\sigma)}$ die konjugierten Zahlen in den Körpern von $(a^0), a', \dots, a^{(n-m-\sigma)}$.

Wir können jetzt die in 14. gewonnenen Sätze über die Substitutionenketten mit n Variablen zu algebraischen Zahlen n^{ten} Grades in der Weise heranziehen, daß wir die Zahl n dort durch den Wert $n - m$ hier ersetzen. Das in (49) liegende Resultat zeigt alsdann, daß man im gegenwärtigen Falle zur beliebigen Zahl r stets $n - m$ ganzzahlige Systeme $y_1 = y_1^{(k)}, \dots, y_{n-m} = y_{n-m}^{(k)} (k = 1, \dots, n - m)$ mit von Null verschiedener Determinante finden

kann, so daß alle Zahlen $y_h^{(k)}$ darin dem Betrage nach $\leq r$ sind, und daß für jedes einzelne dieser $n - m$ Systeme

$$|y_1 + ay_2 + \dots + a^{n-m-1}y_{n-m}| \leq \bar{B}r^{-\frac{n-m-\sigma}{\sigma}}$$

ausfällt, wo $\bar{B}$ eine gewisse nur von a , nicht von r abhängende Konstante ist. Setzt man zu jedem dieser Systeme $y_1, \dots, y_{n-m}$ dann $x_1 = y_1, \dots, x_{n-m} = y_{n-m}$, $x_{n-m+1} = 0, \dots, x_n = 0$, so bekommt man $n - m$ ganzzahlige Systeme $x_1, \dots, x_n$, die von den oben erwähnten speziellen m solchen Systemen unabhängig sind. Danach wird man, sowie $r \geq g^*$ ist, außer $|\alpha_1| = \dots = |\alpha_m| = 0$ weiter stets

$$(60) \quad 0 < |\alpha_{m+1}| \leq \dots \leq |\alpha_n| \leq \bar{B}r^{-\frac{n-m-\sigma}{\sigma}}$$

haben.

Es sei jetzt zuvörderst $n - m > \sigma$. Alsdann erkennt man aus (60), daß die Beträge $|\alpha_{m+1}|, \dots, |\alpha_n|$ im Verlauf der Kette unter jede Grenze sinken, ohne jemals Null zu werden; die Kette bricht also jedenfalls nicht ab. Man hat ferner stets

$$(61) \quad |\alpha_k^{(j)}| \leq \bar{C}r \quad \left(\begin{matrix} k=m+1, \dots, n, \\ j=1, \dots, n-m-\sigma \end{matrix} \right),$$

wenn $\bar{C}$ der größte unter den Werten $1 + |a^{(j)}| + \dots + |a^{(j)}|^{n-1}$ ($j=1, \dots, n-m-\sigma$) ist. Andererseits ist $g_0 a$ eine ganze algebraische Zahl, desgleichen daher jede Zahl $g_0^{n-1} \alpha_k$ ($k=m+1, \dots, n$), und da diese Zahlen von Null verschieden sind, müssen mithin ihre Normen im Körper von a dem Betrage nach ≥ 1 sein; man hat also

$$(62) \quad g_0^{(n-m)(n-1)} |\alpha_k|^\sigma |\alpha'_k \dots \alpha_k^{(n-m-\sigma)}| \geq 1 \quad (k=m+1, \dots, n).$$

Führt man hierin für $n - m - \sigma$ der $n - m - \sigma + 1$ absoluten Beträge $|\alpha_k|, |\alpha'_k|, \dots, |\alpha_k^{(n-m-\sigma)}|$ die in (60) bzw. (61) angewiesenen oberen Grenzen ein, so erhält man für den noch übrigen dieser Beträge eine untere Grenze; man findet

$$(63) \quad |\alpha_k| \geq \bar{b}r^{-\frac{n-m-\sigma}{\sigma}}, |\alpha_k^{(j)}| \geq \bar{c}r \quad \left(\begin{matrix} k=m+1, \dots, n, \\ j=1, \dots, n-m-\sigma \end{matrix} \right)$$

mit gewissen von r nicht abhängenden Konstanten $\bar{b}$ und $\bar{c}$, und daraus folgt dann

$$(64) \quad \left| \frac{\alpha_k}{\alpha_h} \right| \left(= \left| \frac{\alpha_k^0}{\alpha_h^0} \right| \right) \leq \frac{\bar{B}}{\bar{b}}, \quad \left| \frac{\alpha_k^{(j)}}{\alpha_h^{(j)}} \right| \leq \frac{\bar{C}}{\bar{c}} \quad \left(\begin{matrix} h, k=m+1, \dots, n; h \neq k. \\ j=1, \dots, n-m-\sigma \end{matrix} \right).$$

Andererseits findet man aus (60) und (61) die von r nicht abhängende GröÙe $g_0^{(n-m)(n-1)} \bar{B}^\sigma \bar{C}^{n-m-\sigma}$ als obere Grenze für die Beträge der Normen von $g_0^{n-1} \alpha_k$ ($k=m+1, \dots, n$). Aus (64) und aus diesem letzten Umstände schließt man endlich durch die entsprechende Überlegung wie bei (57), daß für die Verhältnisse der Koeffizienten $\alpha_{m+1}, \dots, \alpha_n$ in χ (die

Zahl $r \geq g^*$ angenommen) nur eine endliche Anzahl von algebraischen Zahlen in Betracht kommen. Danach sind in der Tat sämtliche Formen $\chi_1, \chi_2, \dots$ der Kette aus einer endlichen Anzahl unter ihnen durch Multiplikation mit Faktoren abzuleiten. Damit ist der Punkt 3^o des Kriteriums erwiesen. Indem man noch den Umstand heranzieht, daß auch im gegenwärtigen Falle nach 9. die Determinante jeder Substitution der Kette dem Betrage nach $\leq n!$ ist, kann man wieder hinzufügen, daß sich unter den Formen $\chi_1, \chi_2, \dots$ eine endliche Anzahl hervorheben läßt, aus denen alle diese Formen durch Multiplikation mit solchen Faktoren entstehen, welche *Einheiten* in dem Zahlenkörper von a sind.

Es sei endlich $n - m = \sigma$. Ist $\sigma = 2$, also a komplex, so bleiben in χ für $r \geq g^*$ allein die Koeffizienten α_{n-1}, α_n von Null verschieden. In den quadratischen Gleichungen

$$(t - g_0^{n-1}\alpha_{n-1})(t - g_0^{n-1}\alpha_n^0) = 0, \quad (t - g_0^{n-1}\alpha_n)(t - g_0^{n-1}\alpha_n^0) = 0$$

sind die Koeffizienten ganze rationale Zahlen und hat man durch (60) von r unabhängige obere Grenzen für die Beträge derselben. Danach kommen für α_{n-1}, α_n nur eine endliche Anzahl von Werten in Betracht. Da nun nach einer Bemerkung bei (30) alle Formen χ_x der Kette verschieden ausfallen, muß danach die Kette nach einer endlichen Anzahl von Gliedern abbrechen. — Ist $\sigma = 1$, also a reell und rational, so bleibt für $r \geq g^*$ nur α_n von Null verschieden, dabei ist $g_0^{n-1}\alpha_n$ eine ganze rationale Zahl und liegt zufolge (60) dem Betrage nach unterhalb einer gewissen von r unabhängigen Grenze. Die Kette bricht daher auch hier nach einer endlichen Anzahl von Gliedern ab. Damit ist auch der letzte Punkt des in 10. aufgestellten Satzes bewiesen.

Zürich, den 9. Februar 1899.

XV.

Zur Theorie der Einheiten in den algebraischen Zahlkörpern.

(Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen.

Mathematisch-physikalische Klasse. 1900. S. 90—93.)

(Vorgelegt von D. Hilbert in der Sitzung vom 3. März 1900.)

In der Theorie der algebraischen Zahlen ist die Frage von Interesse, ob in einem beliebigen Galoisschen Körper stets eine Einheit existiert, welche unter ihren konjugierten Zahlen ein vollständiges System unabhängiger Einheiten des Körpers darbietet. Durch die folgenden Betrachtungen wird diese Frage in bejahendem Sinne entschieden.

1. Ich benutze dabei hauptsächlich den nachstehenden Satz über das Nichtverschwinden einer gewissen Determinante:

Hilfssatz. Sind in einer m-reihigen Determinante

$$A = |a_{hk}| \quad (h, k = 1, 2, \dots, m)$$

von m^2 reellen Größen alle Glieder a_{hk} außerhalb der Hauptdiagonale (also für $h \neq k$) < 0 und sind ferner die Summen

$$a_{h1} + a_{h2} + \dots + a_{hm} = s_h \quad (h = 1, 2, \dots, m)$$

der Glieder in den einzelnen Horizontalreihen durchweg > 0 , so ist die Determinante stets > 0 .

Beweis. Für $m = 1$ ist der Satz selbstverständlich. Um ihn für einen bestimmten Wert $m > 1$ zu erweisen, nehmen wir an, daß er für alle kleineren Werte des m bereits sichergestellt sei. Wir setzen $a_{hh} = s_h + a_{hh}^*$ ($h = 1, 2, \dots, m$) und entwickeln die Determinante A nach den m Größen $s_1, s_2, \dots, s_m$. Das von $s_1, s_2, \dots, s_m$ freie Glied dieser Entwicklung wird eine m -reihige Determinante A^* , bei welcher in jeder Horizontalreihe die Summe aller Glieder Null ist und welche daher den Wert Null hat. Weiter wird der Koeffizient eines Produkts $s_{k_1} s_{k_2} \dots s_{k_l}$, wenn $l \geq 1$ und $< m$ ist, eine Determinante von $m - l$, also weniger als m Reihen, in welcher alle Glieder außerhalb der Hauptdiagonale negativ sind und in jeder Horizontalreihe die Summe der Glieder größer als die Summe der Glieder der entsprechenden Horizontalreihe von A^* , also gewiß > 0 ist. Auf Grund

der von uns gemachten Voraussetzung erweist sich daher jeder dieser Koeffizienten > 0 . Außerdem erscheint in jener Entwicklung noch das Glied $s_1 s_2 \dots s_m$. Nach dieser Zusammensetzung der Determinante A aus lauter positiven Termen ist sie offenbar > 0 .

2. Es sei jetzt θ eine *algebraische Zahl* n^{ten} Grades und unter den n verschiedenen konjugierten algebraischen Zahlen $\theta_1, \theta_2, \dots, \theta_n$, von denen θ ein Wert ist, seien r reelle und $\frac{1}{2}(n-r)$ Paare von konjugiert imaginären Zahlen vorhanden. Von dem Falle $n=2, r=0$ wollen wir absehen. Wir setzen $r + \frac{1}{2}(n-r) = m+1$, und wir wollen annehmen, daß in der Reihe der $m+1$ ersten Zahlen $\theta_1, \theta_2, \dots, \theta_{m+1}$ die sämtlichen reellen jener Zahlen und ferner je eine Zahl aus jedem der Paare konjugiert imaginärer Zahlen sich befinden; zudem möge die Zahl θ selbst unter diesen $m+1$ Zahlen vorhanden sein. Ist ε eine Einheit in dem algebraischen Körper von θ , so wollen wir mit $l_h(\varepsilon)$ für $h=1, 2, \dots, m+1$ den reellen Teil des Logarithmus der zu ε konjugierten Zahl im Körper von θ_h oder das Doppelte dieses reellen Teils verstehen, je nachdem die Zahl θ_h reell oder imaginär ist. Da die Norm einer Einheit gleich ± 1 ist, gilt dann stets:

$$(1) \quad l_1(\varepsilon) + l_2(\varepsilon) + \dots + l_{m+1}(\varepsilon) = 0.$$

Wie Dirichlet gezeigt hat, kann man in dem Körper von θ stets eine Einheit derart bestimmen, daß von ihren konjugierten Zahlen in den Körpern von $\theta_1, \theta_2, \dots, \theta_{m+1}$ alle bis auf eine Zahl in einem beliebig angenommenen dieser Körper absolute Beträge < 1 haben. Es sei in solcher Weise $\varepsilon^{(h)}$ für $h=1, 2, \dots, m$ eine Einheit, für welche $l_1(\varepsilon^{(h)}), \dots, l_{h-1}(\varepsilon^{(h)}), l_{h+1}(\varepsilon^{(h)}), \dots, l_{m+1}(\varepsilon^{(h)})$, also sämtliche Werte $l_k(\varepsilon^{(h)})$ für $k \neq h$ negativ ausfallen. Dann ist mit Rücksicht auf (1):

$$l_1(\varepsilon^{(h)}) + \dots + l_h(\varepsilon^{(h)}) + \dots + l_m(\varepsilon^{(h)}) > 0.$$

Die Determinante

$$|l_k(\varepsilon^{(h)})| \quad (h, k = 1, 2, \dots, m)$$

trägt danach diesen Charakter, daß in ihr jeder Koeffizient außerhalb der Hauptdiagonale negativ ist und die Summe der Glieder in jeder Horizontalreihe positiv ist. Dem Hilfssatze in 1. zufolge ist daher diese Determinante > 0 . Somit bilden die hier charakterisierten Einheiten $\varepsilon^{(1)}, \varepsilon^{(2)}, \dots, \varepsilon^{(m)}$ ein *vollständiges* System von unabhängigen Einheiten im Körper von θ .

Bei der Methode, welche Dirichlet zur Herstellung eines vollständigen Systems von unabhängigen Einheiten in einem Zahlkörper gegeben hat, werden die Einheiten des Systems sukzessive bestimmt, wobei zur Ermittlung einer weiteren Einheit, so lange das System noch nicht vollständig ist, gewisse Determinanten aus den Logarithmen der früher bestimmten Einheiten und ihrer konjugierten Zahlen bekannt sein müssen.

Nach dem hier Auseinandergesetzten ist es dagegen stets möglich, Einheiten in der erforderlichen Anzahl gesondert, jede für sich, zu bestimmen mit dem Erfolge, daß sie zusammen ein vollständiges System unabhängiger Einheiten ergeben.

3. Es sei jetzt der Körper von θ ein Galois'scher Körper, so daß sich jede der Zahlen $\theta_1, \theta_2, \dots, \theta_n$ als eine rationale Funktion mit rationalen Koeffizienten von jeder anderen dieser Zahlen darstellen läßt. Es sei in solcher Weise

$$\theta_h = R_h(\theta), \quad \theta = S_h(\theta_h) \quad (h = 1, 2, \dots, n),$$

wo R_h, S_h Zeichen für rationale Funktion mit rationalen Koeffizienten sind, so gilt $S_h(R_h(\theta)) = \theta$ und infolgedessen auch allgemein $S_h(R_h(\theta_k)) = \theta_k$ für jeden Index $k = 1, 2, \dots, n$. Je nachdem θ reell oder imaginär ist, sind auch die konjugierten Zahlen alle reell oder alle imaginär, und hat man also $m + 1 = n$ oder $= \frac{1}{2}n$.

Wir bestimmen im Körper von θ eine Einheit $\varepsilon = f(\theta)$, wo $f(\theta)$ eine rationale Funktion mit rationalen Koeffizienten von θ bedeute, so daß von den $m + 1$ konjugierten Zahlen $f(\theta_1), f(\theta_2), \dots, f(\theta_{m+1})$ alle mit Ausnahme der einen Zahl $f(\theta)$ absolute Beträge < 1 haben. Sodann sei ε_h für $h = 1, 2, \dots, m + 1$ die konjugierte Zahl zu ε in dem Körper von θ_h , also $\varepsilon_h = f(\theta_h) = f(R_h(\theta))$; dabei ist ε_h jedesmal selbst eine Zahl im Körper von θ und sind deren konjugierte Zahlen in den Körpern von $\theta_1, \theta_2, \dots, \theta_{m+1}$ bezüglich

$$(2) \quad f(R_h(\theta_1)), f(R_h(\theta_2)), \dots, f(R_h(\theta_{m+1})).$$

Nun hat man bei beliebigen Werten h, k aus der Reihe $1, 2, \dots, m + 1$ stets $R_h(\theta_k) = R_h(R_k(\theta)) = \theta_g$, wo g einen der Indizes $1, 2, \dots, n$ bedeutet. Dabei kann nicht für zwei verschiedene Indizes k bei demselben Index h hier derselbe Wert θ_g und können auch nicht konjugiert imaginäre Werte θ_g resultieren, da aus dieser Relation umgekehrt $\theta_k = S_h(\theta_g)$ folgt und unter den Zahlen $\theta_1, \theta_2, \dots, \theta_{m+1}$ keine zwei gleich oder konjugiert imaginär sind. Danach sind die absoluten Beträge der Größen in (2) abgesehen von der Reihenfolge identisch mit den Beträgen der Größen $f(\theta_1), f(\theta_2), \dots, f(\theta_{m+1})$, es ist also einer jener Beträge > 1 und die m anderen sind sämtlich < 1 , und zwar ist für denjenigen Index k der Betrag $|f(R_h(\theta_k))| > 1$, für den $R_h(\theta_k)$ gleich θ bezüglich gleich der konjugiert imaginären Zahl, also θ_k gleich $S_h(\theta)$ bezüglich gleich der konjugiert imaginären Zahl ist. Für verschiedene Werte h der Reihe $1, 2, \dots, m + 1$ wird der hier in Betracht kommende Index k aus der Reihe $1, 2, \dots, m + 1$ jedesmal ein anderer sein. Nach den Auseinandersetzungen in 2. bilden nunmehr irgend m der Zahlen $\varepsilon_1, \varepsilon_2, \dots, \varepsilon_{m+1}$ ein vollständiges System von unabhängigen Ein-

heiten in dem Galoisschen Körper von θ . Auf diese Weise resultiert der Satz:

In einem Galoisschen Körper kann man stets eine solche Einheit angeben, daß eine jede Einheit dieses Körpers ein Produkt aus einer Einheitswurzel und aus Potenzen dieser Einheit und ihrer konjugierten Einheiten mit rationalen Exponenten ist.

Da ein jeder algebraischer Körper ein Unterkörper eines Galoisschen Körpers ist und seine Einheiten dann zugleich als Einheiten des Galoisschen Körpers auftreten, ist hierdurch zugleich ein bemerkenswerter Satz über die Einheiten in einem beliebigen algebraischen Körper gewonnen.

Zürich, den 23. Februar 1900.

XVI.

Über die Annäherung an eine reelle GröÙe durch rationale Zahlen.

(Mathematische Annalen, Band 54, S. 91—124.)

Obwohl die Theorie der Annäherung an eine reelle GröÙe mit Hilfe von Kettenbrüchen seit Euler und Lagrange noch mannigfache Behandlung erfahren hat, scheint einer der interessantesten Sätze auf diesem Gebiete bisher nicht bemerkt worden zu sein. Nämlich unter den verschiedenen möglichen Kettenbruchentwicklungen für eine reelle GröÙe a , wobei die Teilzähler ± 1 und die Teilnenner positive ganze Zahlen sind, gibt es eine bestimmte Art der Entwicklung (und zwar die am besten konvergierende), für welche die sämtlichen Näherungsbrüche $\frac{x}{y}$ sich von vornherein in einfachster Weise charakterisieren lassen: Als Zähler und Nenner der einzelnen Näherungsbrüche erscheinen dabei genau die sämtlichen Paare von ganzen Zahlen x, y , für die $y > 0$ ist, x und y relativ prim sind und dazu die Bedingung

$$|(x - ay)y| < \frac{1}{2}$$

erfüllt ist. Von dem Falle, daß a gleich einer ganzen Zahl $+\frac{1}{2}$ ist, hat man hierbei abzusehen.*)

*) Bekanntlich läßt sich eine nach fallenden Potenzen von z fortschreitende konvergente Reihe

$$f(z) = c_{-m}z^m + c_{-m+1}z^{m-1} + \dots + c_0 + \frac{c_1}{z} + \frac{c_2}{z^2} + \dots$$

in einen Kettenbruch

$$F_0(z) + \frac{1}{F_1(z)} + \frac{1}{F_2(z)} + \dots$$

umwandeln, so daß $F_0(z)$ eine ganze rationale Funktion von z und $F_1(z), F_2(z), \dots$ ganze rationale Funktionen von z mindestens vom Grade 1 sind. Dabei gilt der Satz: Ein Quotient $\frac{P(z)}{Q(z)}$ zweier relativ primer ganzer rationaler Funktionen von z ist immer dann und nur dann Näherungsbruch dieses Kettenbruchs, wenn die Entwicklung von

$$(P(z) - f(z)Q(z))Q(z)$$

nach fallenden Potenzen von z mit einer Potenz von z , deren Exponent negativ ist,

Auf die betreffende noch durch weitere bemerkenswerte Eigenschaften ausgezeichnete Kettenbruchentwicklung habe ich bereits an anderer Stelle*) hingewiesen, ohne jedoch damals wahrzunehmen, daß die eben erwähnte Ungleichung für die Näherungsbrüche diese Entwicklung bereits vollständig charakterisiert.

Im folgenden gebe ich eine auf geometrischen Betrachtungen gegründete und dadurch sehr anschauliche *Theorie des Systems zweier linearer Formen* $\alpha x + \beta y$, $\gamma x + \delta y$ mit beliebigen reellen Koeffizienten und mit ganzzahligen Unbestimmten. Eine Anwendung der dabei zutage tretenden Resultate auf die spezielleren Ausdrücke $x - ay$, y liefert dann insbesondere jene Sätze über die Annäherung an eine GröÙe a .

§ 1.

Satz I. Sind $\xi = \alpha x + \beta y$, $\eta = \gamma x + \delta y$ zwei lineare Formen mit beliebigen reellen Koeffizienten $\alpha, \beta, \gamma, \delta$ und einer Determinante $\alpha\delta - \beta\gamma = 1$, so gibt es stets ganze Zahlen x, y , die nicht beide Null sind und für welche

$$|\xi\eta| \leq \frac{1}{2}$$

ausfällt.

Wird auf die Form $\xi\eta$ eine ganzzahlige Substitution $x = pX + p'Y$, $y = qX + q'Y$ mit einer Determinante ± 1 angewandt, so nennen wir $\xi\eta$ der transformierten Form in den neuen Variablen X, Y äquivalent. Zugleich mit dem Satze I beweisen wir den folgenden

Zusatz. Ist $\xi\eta$ weder mit der Form XY noch mit der Form $\frac{1}{2}(X^2 - Y^2)$ äquivalent, so gibt es stets ganze Zahlen x, y , wofür $\xi \neq 0$, $\eta \neq 0$ und $|\xi\eta| < \frac{1}{2}$ ausfällt.

Beweis. Wir deuten x, y als irgendwelche Parallelkoordinaten in einer Ebene, wobei noch für jede der zwei Koordinaten die Entfernung Eins parallel ihrer Achse in willkürlicher Weise angenommen sein kann. Es sei O der Nullpunkt ($x = 0, y = 0$). Bedeutet A einen von O verschiedenen Punkt $x = p, y = q$, so soll der zu ihm in bezug auf O symmetrische Punkt $x = -p, y = -q$ jedesmal mit A_0 bezeichnet werden.

Die Gesamtheit derjenigen Punkte, für welche x wie y ganze Zahlen sind, heiÙe das *Zahlengitter*, und die einzelnen Punkte daraus sollen *Gitterpunkte* heiÙen.

beginnt. Der Satz, den ich im Texte angebe, stellt das wohl von manchem Mathematiker vermiÙte Analogon in der GröÙenlehre zu diesem Satze der Funktionenlehre vor.

*) Annales de l'École Normale supérieure, 3^e série, T. XIII; 1896. Diese Ges. Abhandlungen, Bd. I, S. 278.

Sind ϱ, σ positive Parameter, so bilden die vier Punkte R, R_0, S, S_0 , für welche $\xi = \varrho, \eta = 0$; $\xi = -\varrho, \eta = 0$; $\xi = 0, \eta = \sigma$; $\xi = 0, \eta = -\sigma$ ist, die Ecken für ein Parallelogramm mit O als Mittelpunkt und mit den Linien $\xi = 0, \eta = 0$ als *Diagonalen*. Ein solches Parallelogramm werde mit $\mathfrak{P}(\varrho, \sigma)$ bezeichnet. Seine vier Seiten besitzen die Gleichungen $\pm \frac{\xi}{\varrho} \pm \frac{\eta}{\sigma} = 1$, wo für die zwei Vorzeichen alle vier Kombinationen $+, +; -, +; -, -; +, -$ in Betracht kommen, und der Bereich von $\mathfrak{P}(\varrho, \sigma)$ wird daher durch

$$\left| \frac{\xi}{\varrho} \right| + \left| \frac{\eta}{\sigma} \right| \leq 1$$

dargestellt.

Wir können nun von so kleinen Werten für ϱ und σ ausgehen, daß das zugehörige Parallelogramm $\mathfrak{P}(\varrho, \sigma)$ jedenfalls keinen Gitterpunkt außer dem Nullpunkte O in sich faßt. Dann lassen wir ϱ und σ *unter Festhaltung der Größe ihres Verhältnisses* $\varrho:\sigma$ wachsen. Dabei dehnt sich $\mathfrak{P}(\varrho, \sigma)$ nach allen Richtungen von O aus in gleichem Maße aus. Wir müssen daher schließlich zu gewissen Werten ϱ, σ kommen, wobei $\mathfrak{P}(\varrho, \sigma)$ auf seiner *Begrenzung* weitere Gitterpunkte aufnimmt, während immer noch O der einzige Gitterpunkt im *Inneren* von $\mathfrak{P}(\varrho, \sigma)$ ist. Es sei bei diesen Werten ϱ, σ , bei denen wir nun verweilen, $A (x=p, y=q)$ ein Gitterpunkt auf der Begrenzung von $\mathfrak{P}(\varrho, \sigma)$. Da zugleich mit A auch der Gitterpunkt $A_0 (x=-p, y=-q)$ auf der Begrenzung von $\mathfrak{P}(\varrho, \sigma)$ auftritt, so können wir annehmen, für A sei $\eta > 0$ oder $\eta = 0, \xi > 0$. Wir setzen für A : $\xi = \varepsilon \lambda, \eta = \mu$, so daß $\mu \geq 0, \lambda \geq 0, \varepsilon = \pm 1$ ist.

Die ganzen Zahlen p, q haben gewiß keinen gemeinsamen Teiler > 1 , weil die Strecke OA keinen Gitterpunkt zwischen O und A enthält. Wir bestimmen zwei ganze Zahlen r, s irgendwie so, daß $ps - qr = \varepsilon$ ist, und setzen

$$x = p\bar{X} + rY, \quad y = q\bar{X} + sY.$$

Dabei werde $\varepsilon \xi = \lambda \bar{X} + \bar{\lambda} Y, \eta = \mu \bar{X} + \bar{\mu} Y$. Dann haben $\varepsilon \xi, \eta$ in $\bar{X}, Y$ die Determinante $\varepsilon \varepsilon = 1$ und folgt

$$(1) \quad Y = \lambda \eta - \mu \varepsilon \xi.$$

Die Gitterpunkte x, y werden genau die Punkte mit ganzzahligen Bestimmungsstücken $\bar{X}, Y$. Diese Punkte ergeben auf der Linie $Y = 0$, d. i. der Geraden durch O und A die unendliche Punktreihe zu den Werten $\bar{X} = \dots - 2, -1, 0, 1, 2, \dots$, wobei $\bar{X} = 0$ den Nullpunkt, $\bar{X} = 1$ den Punkt A liefert und alle diese Punkte sich in einem konstanten Abstände $= OA$ folgen. Sodann bilden sie auf einer jeden der zu $Y = 0$ parallelen Geraden $Y = 1, Y = -1, Y = 2, Y = -2, \dots$ jedesmal eine äquidistante Punktreihe mit dem gleichen konstanten Abstände $= OA$ zwischen benachbarten Punkten. Von diesen sämtlichen

einander parallelen Geraden sind $Y = 1$ und $Y = -1$ die zwei an $Y = 0$ nächstgelegenen.

1. Wir nehmen zunächst an, daß A *nicht* eine *Ecke* von $\mathfrak{P}(\rho, \sigma)$, also $\lambda > 0$, $\mu > 0$ ist. *) Es sei F der Punkt $\xi = -\varepsilon\lambda$, $\eta = \mu$. Das Parallelogramm mit den Ecken A , F , A_0 , F_0 ist durch $-\lambda \leq \xi \leq \lambda$, $-\mu \leq \eta \leq \mu$ definiert und befindet sich, von den Ecken abgesehen, ganz im Inneren des Parallelogramms $\mathfrak{P}(\rho, \sigma)$.

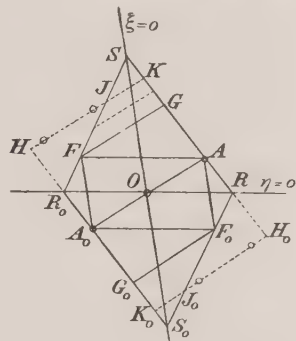


Fig. 1.

Nehmen wir weiter an, daß A auch *nicht* *Mitte* einer Seite von $\mathfrak{P}(\rho, \sigma)$ ist. Die zu A_0OA parallele Gerade durch F trifft dann die Begrenzung von $\mathfrak{P}(\rho, \sigma)$ außer in F in einem zweiten Punkte G , so daß die Strecke FG offenbar *größer* als OA ist. Ebenso würde jede parallele Gerade zu A_0OA , die in dem Streifen zwischen A_0OA und FG verläuft, eine Strecke $> OA$ ganz im Inneren von $\mathfrak{P}(\rho, \sigma)$ liegen haben. Da nun im Inneren von $\mathfrak{P}(\rho, \sigma)$ sich kein Gitterpunkt außer O befindet, welcher Punkt auf $Y = 0$ liegt, müssen danach die Geraden $Y = \pm 1$ außerhalb des durch die Parallelen FG und F_0G_0 begrenzten Streifens verlaufen, d. h. für den Punkt F muß jedenfalls $|Y| < 1$ sein. Nun hat man für F : $Y = 2\lambda\mu$, also folgt

$$2\lambda\mu < 1.$$

Danach ist der Gitterpunkt A hier von der im Satze I und dem Zusatze geforderten Beschaffenheit.

2. Ist A *Mitte* einer Seite von $\mathfrak{P}(\rho, \sigma)$, also $\lambda = \frac{1}{2}\rho$, $\mu = \frac{1}{2}\sigma$, so ist A_0OA , also die Gerade $Y = 0$ parallel einer Seite von $\mathfrak{P}(\rho, \sigma)$. Jede Parallele zu A_0OA , welche ins Innere von $\mathfrak{P}(\rho, \sigma)$ eintritt, hat dann mit $\mathfrak{P}(\rho, \sigma)$ eine Strecke $= A_0OA > OA$ gemein. Die Geraden $Y = \pm 1$ können daher nicht ins Innere von $\mathfrak{P}(\rho, \sigma)$ eintreten, $\mathfrak{P}(\rho, \sigma)$ liegt also ganz in dem Streifen $-1 \leq Y \leq 1$. Insbesondere gilt danach für den Punkt F : $Y \leq 1$, d. h. man hat

$$2\lambda\mu \leq 1.$$

Der Punkt A hat also wieder die im Satze I verlangte Beschaffenheit.

Das Gleichheitszeichen in dieser Ungleichung, (dessen Eintreten zugleich $\rho\sigma = 2$ bedeuten würde), hat dann statt, wenn eine Seite von

*) Wie die Buchstaben R und R_0 in der Figur eingetragen sind, ist darin $\varepsilon = 1$ für den Punkt A angenommen. Die Gitterpunkte sind hier und in den weiteren Figuren durch kleine Kreise angedeutet.

$\mathfrak{P}(\varrho, \sigma)$ auf die Gerade $Y = 1$ fällt. Da diese Seite an Länge $= 2OA$ ist, so enthält sie alsdann entweder *innerhalb* jeder ihrer durch F getrennten Hälften je einen Gitterpunkt, in welchem Falle für diese Gitterpunkte nach dem bereits in 1. Ausgeführten sich $\xi \neq 0$, $\eta \neq 0$, $|\xi\eta| < \frac{1}{2}$ herausstellt, oder aber: sind auf ihr sowohl beide Ecken wie die Mitte F Gitterpunkte. In diesem zweiten Falle wären in $\mathfrak{P}(\varrho, \sigma)$ alle vier Mitten der Seiten Gitterpunkte. Wir können dann A als die Mitte von RS , d. h. $\varepsilon = 1$ voraussetzen. Ist für F hier $Y = 1$, $\bar{X} = g$ und setzt man $\bar{X} = X + gY$, so sind $A \left(\xi = \frac{1}{2}\varrho, \eta = \frac{1}{2}\sigma \right)$ und $F \left(\xi = -\frac{1}{2}\varrho, \eta = \frac{1}{2}\sigma \right)$ durch $X = 1, Y = 0$ und $X = 0, Y = 1$ bestimmt, und hat man daher

$$\xi = \frac{1}{2}\varrho(X - Y), \quad \eta = \frac{1}{2}\sigma(X + Y), \quad \varrho\sigma = 2, \quad \xi\eta = \frac{1}{2}(X^2 - Y^2).$$

3. Endlich nehmen wir an, daß der Gitterpunkt A eine *Ecke* von $\mathfrak{P}(\varrho, \sigma)$, also dafür $\xi = 0$ oder $\eta = 0$ sei; dann ist selbstverständlich für diesen Punkt $|\xi\eta| = 0 < \frac{1}{2}$. In jedem Falle entspricht somit der Gitterpunkt A den Bedingungen des Satzes I.

Um auch noch den Zusatz unter den letzten Umständen als richtig zu erkennen, stellen wir uns A etwa als die Ecke $R(\xi = \varrho, \eta = 0)$ von $\mathfrak{P}(\varrho, \sigma)$ vor. Dann ist nach (1.): $Y = \varrho\eta$. Nunmehr halten wir ϱ fest und denken uns den Parameter σ als veränderlich. Dabei bleibt die Diagonale $R_0R(A_0OA)$ von $\mathfrak{P}(\varrho, \sigma)$ auf $Y = 0$ fest, während sich die Endpunkte S_0, S der anderen Diagonale auf der Linie $\xi = 0$ verschieben. Bei hinreichend kleinem σ liegt $\mathfrak{P}(\varrho, \sigma)$ ganz im Bereiche $-1 < Y < 1$, enthält dann also gewiß keinen Gitterpunkt außer A_0, O, A . Wird $\sigma = \frac{1}{\varrho}$, so erreicht $\mathfrak{P}(\varrho, \sigma)$ die Gerade $Y = 1$ mit der Ecke $S(\xi = 0, \eta = \sigma)$. Fällt dabei diese Ecke zugleich in einen Gitterpunkt, so sei für ihn $Y = 1$, $\bar{X} = g$. Setzt man alsdann $\bar{X} = X + gY$, so gilt für S : $\xi = 0$, $\eta = \sigma$; $X = 0$, $Y = 1$, für R : $\xi = \varrho$, $\eta = 0$; $X = 1$, $Y = 0$, also hat man in diesem Falle:

$$\xi = \varrho X, \quad \eta = \sigma Y, \quad \varrho\sigma = 1, \quad \xi\eta = XY.$$

Fällt hingegen der Punkt $\xi = 0, \eta = \frac{1}{\varrho}$ nicht in einen Gitterpunkt, so kann man σ über $\frac{1}{\varrho}$ hinaus wachsen lassen, ohne daß zunächst neue Gitterpunkte in $\mathfrak{P}(\varrho, \sigma)$ eintreten. Dabei wird die Strecke, welche der Bereich von $\mathfrak{P}(\varrho, \sigma)$ aus der Linie $Y = 1$ ausschneidet und deren variable Endpunkte J, K heißen mögen, nach beiden Enden zu immer ausgedehnter und wird sich schließlich einer dieser zwei Fälle ereignen:

Entweder fällt, so lange noch diese Strecke $JK < OA$ ist, einer

ihrer Endpunkte (wie in Fig. 2 der Punkt J) in einen Gitterpunkt A' auf der Geraden $Y=1$. Dabei reicht dann wegen $JK < \frac{1}{2} A_0 O A$ das Parallelogramm $\mathfrak{P}(\varrho, \sigma)$ noch nicht an $Y=2$ heran, enthält also gewiß keinen Gitterpunkt außer O, A, A_0, A', A'_0 , und zugleich liegt der Punkt A' auf $Y=1$ näher an S (wofür $Y < 2$ ist), als an dem anderen Endpunkte (A_0 in Fig. 2) der durch ihn gehenden Seite von $\mathfrak{P}(\varrho, \sigma)$, also ist A' keinesfalls Mitte dieser Seite. In diesem Falle gilt für A' nach den früheren Bemerkungen $\xi \neq 0, \eta \neq 0, |\xi\eta| < \frac{1}{2}$.

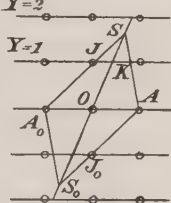


Fig. 2.

Der andere mögliche Fall ist, daß erst, wenn die Strecke JK bis zur Länge OA gewachsen ist, ihre beiden Endpunkte in Gitterpunkte zu liegen kommen. In dieser Endlage stößt dann $\mathfrak{P}(\varrho, \sigma)$ gerade mit der Ecke S auf $Y=2$, so daß auch hier noch $\mathfrak{P}(\varrho, \sigma)$ keinen Gitterpunkt außer O im Inneren enthält, aber in den Mitten aller vier Seiten Gitterpunkte aufweist. In diesem Falle erweist sich, wie schon oben unter 2. ausgeführt ist, $\xi\eta$ als äquivalent mit $\frac{1}{2}(X^2 - Y^2)$. Werden alle erörterten Einzelheiten zusammengefaßt, so tritt die Richtigkeit des zum Satze I gemachten Zusatzes hervor. —

Den Beweis der Ungleichung $2\lambda\mu \leq 1$ für den Gitterpunkt A und damit des Satzes I können wir auch durch folgende, auf dem Begriffe des *Flächeninhalts* beruhende Betrachtung erhalten, die noch zu einem weiter reichenden Resultate führt.

Da wir im Hinblick auf die festzustellende Relation $|\xi\eta| \leq \frac{1}{2}$ die Rolle der Formen ξ, η vertauschen, auch ξ durch $-\xi$ ersetzen können, wobei freilich als Wert der Determinante von ξ, η auch -1 zuzulassen ist, so dürfen wir ohne wesentliche Einschränkung annehmen, A falle auf die Seite RS von $\mathfrak{P}(\varrho, \sigma)$ und noch derart, daß $RA \leq AS$ ist, d. h. wir setzen $\varepsilon = 1$ und $\frac{\lambda}{\varrho} \geq \frac{\mu}{\sigma}$ voraus. Indem A auf der Begrenzung von $\mathfrak{P}(\varrho, \sigma)$ liegt, hat man

$$(2) \quad \frac{\lambda}{\varrho} + \frac{\mu}{\sigma} = 1,$$

also $\frac{\lambda}{\varrho} = \frac{1}{2} + \omega, \frac{\mu}{\sigma} = \frac{1}{2} - \omega, \frac{\lambda\mu}{\varrho\sigma} = \frac{1}{4} - \omega^2 \leq \frac{1}{4}$. Nun werden wir beweisen, daß für ein Parallelogramm wie $\mathfrak{P}(\varrho, \sigma)$, welches einen Gitterpunkt A auf der Berandung, im Inneren aber O als einzigen Gitterpunkt enthält, stets

$$(3) \quad \varrho\sigma \leq 2$$

gilt. Daraus folgt dann unmittelbar $\lambda\mu \leq \frac{1}{2}$. Dieser Satz (3) ist aber einschneidender.

Der *Flächeninhalt* von $\mathfrak{P}(\varrho, \sigma)$, d. h. das Doppelintegral $\iint dx dy$, über diesen Bereich erstreckt, ist, da ξ, η die Determinante ± 1 in x, y haben, $= 4 \cdot \frac{1}{2} \varrho \sigma = 2 \varrho \sigma$. Es seien nun H, K (Fig. 1) die Schnittpunkte von $Y = 1$ mit den Geraden $S_0 R_0$ und RS . Wegen (2) haben die Formen $\frac{\xi}{\varrho} + \frac{\eta}{\sigma}$ und Y die Determinante 1 in den Variablen $\bar{X}, Y$ und dann eine Determinante ± 1 in x, y . Also ist der Flächeninhalt des Parallelogramms $K_0 H_0 K H$ gleich 4.

Tritt nun die Strecke HK in das *Innere* von $\mathfrak{P}(\varrho, \sigma)$ ein, so liegt K auf RS zwischen A und S und hat HK mit $\mathfrak{P}(\varrho, \sigma)$ eine Strecke JK gemein, die notwendig $\leq OA$ ist, weil kein Gitterpunkt außer O im Inneren von $\mathfrak{P}(\varrho, \sigma)$ liegt. Wegen $JK \leq OA$ liegt J auf $R_0 S$ und so, daß $R_0 J \geq JS$ ist. Infolgedessen ist der Flächeninhalt des Dreiecks JKS nicht größer als der des Dreiecks JHR_0 . Nun entsteht das Parallelogramm $RSR_0 S_0$ aus dem Parallelogramm $K_0 H_0 K H$, indem von letzterem die Dreiecke $J_0 H_0 R$ und JHR_0 fortgenommen und die Dreiecke $J_0 K_0 S_0$ und JKS von nicht größerem Flächeninhalte hinzugefügt werden. Daraus folgt für die Flächeninhalte der zwei Parallelogramme das Verhältnis $2 \varrho \sigma \leq 4$.

Tritt jedoch die Strecke HK überhaupt nicht in das Innere von $\mathfrak{P}(\varrho, \sigma)$ ein, so ist das Parallelogramm $RSR_0 S_0$ ganz im Parallelogramme $K_0 H_0 K H$ enthalten und daher ebenfalls $2 \varrho \sigma \leq 4$.

§ 2.

1. Es mögen ξ und η dieselbe Bedeutung wie in Satz I haben. Wir wollen jedoch jetzt von vornherein die besonderen Fälle ausschließen, daß die Form $\xi \eta$ der Variablen x, y mit der Form XY oder mit der Form $\frac{1}{2}(X^2 - Y^2)$ der Variablen X, Y äquivalent ist. Von diesen Ausnahmefällen abgesehen, gilt der

Satz II. Sind $x = p, y = q$ zwei relativ prime ganze Zahlen, wofür $|\xi| > 0$ und $|\xi \eta| < \frac{1}{2}$ ausfällt, so kann man stets zwei ganze Zahlen $x = p', y = q'$ finden, so daß $p q' - q p' = \pm 1$ ist und für welche ebenfalls $|\xi \eta| < \frac{1}{2}$ ist, aber $|\xi|$ kleiner ausfällt als für das erstere System.

Beweis. Da wir anstatt p, q ebensogut das System $-p, -q$ zugrunde legen können, nehmen wir an, für $x = p, y = q$ sei $\xi = \varepsilon \lambda, \eta = \mu$ und dabei $\mu > 0, \lambda > 0, \varepsilon = \pm 1$ oder $\mu = 0, \lambda > 0, \varepsilon = 1$. Wir bestimmen zwei ganze Zahlen r, s irgendwie so, daß

$$ps - qr = \varepsilon$$

ist, und setzen

$$x = p\bar{X} + rY, \quad y = q\bar{X} + sY.$$

Alsdann sei

$$\varepsilon \xi = \lambda \bar{X} + \bar{\lambda} Y, \quad \eta = \mu \bar{X} + \bar{\mu} Y,$$

so folgt noch

$$\lambda \bar{\mu} - \mu \bar{\lambda} = 1, \quad Y = \lambda \eta - \mu \varepsilon \xi = \varepsilon (p y - q x).$$

Wir bezeichnen den Gitterpunkt $x = p$, $y = q$ ($\xi = \varepsilon \lambda$, $\eta = \mu$) mit A , ferner, wenn $\mu > 0$ ist, mit F den Punkt $\xi = -\varepsilon \lambda$, $\eta = \mu$. Für F wird $Y = 2\lambda\mu$, d. i. $Y < 1$ auf Grund unserer Voraussetzung über den Gitterpunkt A . Die Geraden $Y = \pm 1$ schließen also das Parallelogramm mit den Ecken A, F, A_0, F_0 ganz zwischen sich ein, ohne es zu treffen, so daß insbesondere F und F_0 gewiß nicht Gitterpunkte sind.

Die Gerade $Y = 1$ trifft nun die Linien $\xi = -\varepsilon \lambda$, $\xi = 0$, $\xi = \varepsilon \lambda$ in drei Punkten H, J, K (Fig. 3), für die $\eta = \frac{1-\lambda\mu}{\lambda}$, $\eta = \frac{1}{\lambda}$, $\eta = \frac{1+\lambda\mu}{\lambda}$ ist, welche GröÙen sämtlich $> \mu$ sind.

Da $HJ = JK = OA$ ist, so werden auf dieser Geraden $Y = 1$ entweder die drei Punkte H, J, K Gitterpunkte sein, oder es liegt ein Gitterpunkt B innerhalb der Strecke HJ und ein Gitterpunkt C innerhalb der Strecke JK . In diesem letzteren Falle sei dann M der Schnittpunkt der Geraden FB (oder A_0B im Falle $\mu = 0$) mit der Geraden $\xi = 0$ und N der Schnittpunkt der Geraden AC mit der Geraden $\xi = 0$.

Wir bezeichnen nun mit $A'(x = p', y = q')$ *erstens*, wenn J ein Gitterpunkt ist, diesen Gitterpunkt, *zweitens*, wenn in J kein Gitterpunkt

fällt und wenn zugleich M näher an O liegt als N , also $OM < ON$ ist, den Gitterpunkt B , *drittens*, wenn in J kein Gitterpunkt fällt und dabei $OM \geq ON$ ist, den Gitterpunkt C . Da A' auf $Y = 1$ liegt, gilt jedesmal $p'q' - qp' = \varepsilon = \pm 1$. Für A' können wir sodann $\xi = \varepsilon \lambda'$, $\eta = \mu'$ setzen, so daß $\lambda' = 0$ im ersten, $\varepsilon' = -\varepsilon$, $\lambda' > 0$ im zweiten, $\varepsilon' = \varepsilon$, $\lambda' > 0$ im dritten Falle ist, ferner in jedem Falle $\lambda' < \lambda$, $\mu' > \mu$. Im ersten Falle denken wir uns noch $\varepsilon' = -\varepsilon$. Wir wollen ferner hier unter $\mathfrak{P}(\varrho, \sigma)$ mit den Ecken RSR_0S_0 speziell dasjenige völlig bestimmte Parallelogramm mit $\xi = 0$, $\eta = 0$ als Diagonalen verstehen, dessen Berandung sowohl den Punkt A , wie den Punkt A' aufnimmt. Die Ecke $S(\xi = 0, \eta = \sigma)$ kommt dabei im ersten Falle in J , im zweiten in M , im dritten in N zu liegen. Wir können nun zeigen, daß in jedem Falle dieses Parallelogramm $\mathfrak{P}(\varrho, \sigma)$

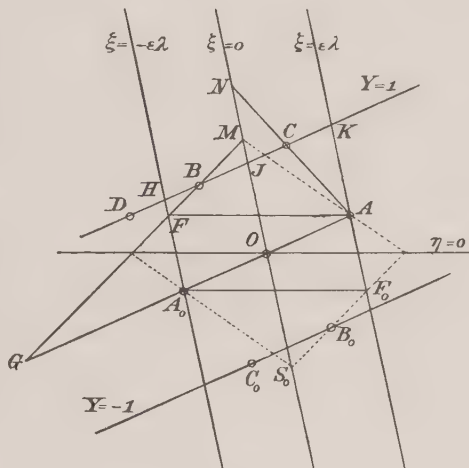


Fig. 3.

keinen Gitterpunkt außer O im Inneren enthält und daß darin auch A' nicht Mitte einer Seite ist. Aus diesen Umständen folgt nach den Betrachtungen in § 1, daß $\lambda'\mu' < \frac{1}{2}$ ist, und, da zudem $\lambda > \lambda' \geq 0$ gilt, wird danach in der Tat der Gitterpunkt p', q' von der im Satze II geforderten Beschaffenheit sein. Wir unterscheiden nun die genannten drei Fälle:

Ist *erstens* J ein Gitterpunkt, $A' = J$, so erreicht das Parallelogramm $\mathfrak{P}(\varrho, \sigma)$ die Gerade $Y=1$ nur im Punkte J . Dieses Parallelogramm enthält daher keinen Gitterpunkt außer O im Inneren und A, A_0, J, J_0 auf der Begrenzung. Dabei werden J, J_0 die Ecken S, S_0 . Da ferner H außerhalb dieses Parallelogramms $\mathfrak{P}(\varrho, \sigma)$ fällt, so liegt wegen $OH = AJ$ der Punkt A von S weiter entfernt als die Mitte $\xi = \frac{\varepsilon\varrho}{2}$, $\eta = \frac{\sigma}{2}$ der Seite, welche A und S enthält. Man hat also $\lambda > \frac{\varrho}{2}$, $\mu < \frac{\sigma}{2}$. Andererseits gilt $\lambda' = 0 < \frac{\varrho}{2}$, $\mu' = \sigma > \frac{\sigma}{2}$. — Zu bemerken ist noch, daß hier jedenfalls $\mu > 0$ ist. Denn wäre $\mu = 0$, also auch A eine Ecke in $\mathfrak{P}(\varrho, \sigma)$, so enthielte dieses Parallelogramm in den vier Ecken Gitterpunkte. Dann wäre nach § 1, 3. die Form $\xi\eta$ der Variablen x, y äquivalent mit der Form XY der Variablen X, Y .

Wir verfolgen jetzt die Annahme, daß J kein Gitterpunkt ist. Es bedeute noch G den Schnittpunkt der Geraden FB mit der Geraden $Y=0$; dieser Punkt liegt im Falle $\mu > 0$ auf der Verlängerung von A_0O über A_0 hinaus, im Falle $\mu = 0$ fällt er mit A_0 zusammen. Aus den zwei ähnlichen Dreiecken GOM und BJM einerseits und aus den zwei ähnlichen Dreiecken OAN und JCN andererseits erhält man die Proportionen

$$(1) \quad \frac{JM}{OJ+JM} = \frac{BJ}{GO}, \quad \frac{JN}{OJ+JN} = \frac{JC}{OA}.$$

Der *zweite* der obigen Fälle, $A' = B$, hat statt, wenn J kein Gitterpunkt ist und dabei $OM < ON$ ist. Der gezeichneten Figur 3 ist speziell dieser Fall zugrunde gelegt. Dabei gibt dann FBM (A_0BM für $\mu = 0$) eine Seite von $\mathfrak{P}(\varrho, \sigma)$ ab. Für den Punkt M gilt hier stets $Y < 2$. Denn entweder hat man $BJ < JC$; wegen $BJ + JC = A_0O$ ist dann $BJ < \frac{1}{2}A_0O \leq \frac{1}{2}GO$ und folgt aus (1): $JM < OJ$. Ist aber $BJ \geq JC$, so ist wegen $BJ + JC = OA$ jetzt $JC \leq \frac{1}{2}OA$ und folgt aus der zweiten Relation in (1): $JN \leq OJ$, und um so mehr gilt dann $JM < OJ$. In jedem Falle ist dann weiter $OM < 2OJ$, d. h. eben $Y < 2$ für den Punkt M .

Das Parallelogramm $\mathfrak{P}(\varrho, \sigma)$ liegt somit hier ganz im Bereiche $-2 < Y < 2$. Von der Geraden $Y=1$ enthält es eine Strecke mit B als einem Endpunkte und mit einem Punkte zwischen J und C als anderem

Endpunkte, von der Geraden $Y=0$ enthält es die Strecke A_0OA . Von Gitterpunkten finden sich also darin allein O im Inneren und A, A_0, B, B_0 auf der Begrenzung. Da ferner $BC=OA$ ist, die Strecke BC bei B ins Innere von $\mathfrak{P}(\rho, \sigma)$ eintritt, aber C außerhalb $\mathfrak{P}(\rho, \sigma)$ fällt, so liegt B an $S(=M)$ näher als die Mitte $\xi = -\frac{\varepsilon\rho}{2}$, $\eta = \frac{\sigma}{2}$ der B und S enthaltenden Seite von $\mathfrak{P}(\rho, \sigma)$, man hat also $\lambda' < \frac{\rho}{2}$, $\mu' > \frac{\sigma}{2}$; und andererseits liegt deshalb der Punkt A von der Ecke $S(=M)$ weiter ab als die Mitte $\xi = \frac{\varepsilon\rho}{2}$, $\eta = \frac{\sigma}{2}$ der A und S enthaltenden Seite von $\mathfrak{P}(\rho, \sigma)$, mithin ist $\lambda > \frac{\rho}{2}$, $\mu < \frac{\sigma}{2}$.

Über die Ermittlung des Gitterpunktes A' in diesen beiden ersten Fällen bemerken wir folgendes: Man hat für A' hier $\xi = -\varepsilon\lambda'$, $\eta = \mu'$ und $0 \leq \lambda' < \lambda$, andererseits $\bar{X}=g$, $Y=1$, wo g eine ganze Zahl ist, und dann $p'=r+gp$, $q'=s+gq$. Nun folgt $-\varepsilon\lambda' = \varepsilon(\bar{\lambda} + g\lambda)$, also soll $0 \leq -\bar{\lambda} - g\lambda < \lambda$ oder $0 \leq -\frac{\bar{\lambda}}{\lambda} - g < 1$ sein. Danach ist g die größte in $-\frac{\bar{\lambda}}{\lambda}$ enthaltene ganze Zahl:

$$g = \left[-\frac{\bar{\lambda}}{\lambda} \right].$$

Die Relation $Y=1$ für A' ist gleichbedeutend mit $\lambda\mu' + \mu\lambda' = 1$. Ist nun $-\frac{\bar{\lambda}}{\lambda}$ genau eine ganze Zahl, so wird $\lambda' = 0$, $A' = J$. Anderenfalls wird $\lambda' > 0$ und ist der Punkt M auf der Geraden FB (A_0B für $\mu = 0$) bestimmt durch $\xi = 0$, $\frac{\eta - \mu'}{\xi + \varepsilon\lambda'} = \frac{\eta - \mu}{\xi + \varepsilon\lambda}$, also für M : $\eta(\lambda - \lambda') = \lambda\mu' - \mu\lambda' = 2\lambda\mu' - 1$, der Punkt N auf der Geraden AC hingegen ist, da C hier den Werten $\xi = -\varepsilon\lambda' + \varepsilon\lambda$, $\eta = \mu' + \mu$ entspricht, bestimmt durch $\xi = 0$, $\frac{\eta - \mu' - \mu}{\xi + \varepsilon\lambda' - \varepsilon\lambda} = \frac{\eta - \mu}{\xi - \varepsilon\lambda}$, also für N : $\eta\lambda' = \lambda\mu' + \mu\lambda' = 1$. Die Relation $OM < ON$ kommt danach auf $\frac{2\lambda\mu' - 1}{\lambda - \lambda'} < \frac{1}{\lambda'}$ d. i., da $\lambda > 0$ ist, auf $2\lambda'\mu' < 1$ hinaus.

Endlich nehmen wir als *dritten* Fall den, daß J kein Gitterpunkt ist und dabei $OM \geq ON$ gilt. Dann hat für A' ($\xi = \varepsilon\lambda'$, $\eta = \mu'$) der Punkt C einzutreten, so daß $\varepsilon' = \varepsilon$ ist. Für B ist dann $\xi = -\varepsilon(\lambda - \lambda')$, $\eta = \mu' - \mu$; es folgt also zunächst $\mu' - \mu > \mu$. Die Relation $OM \geq ON$ kommt hier nach der eben gemachten Ausführung auf $2(\lambda - \lambda')(\mu' - \mu) \geq 1$ hinaus.

Es sei zunächst $OM > ON$. Aus $JM > JN$ folgt nach (1), da $GO \geq OA$ ist, $BJ > JC$. Diese Beziehung ist hier gleichbedeutend mit $\lambda - \lambda' > \lambda'$. Wegen $BJ + JC = OA$ hat man sodann $JC < \frac{1}{2}OA$,

und wegen (1) daher $JN < OJ$. Danach gilt für den Punkt N jedenfalls $Y < 2$. Das Parallelogramm $\mathfrak{P}(\varrho, \sigma)$ reicht also nicht an $Y = 2$ heran. Von der Geraden $Y = 1$ enthält es eine Strecke mit einem Punkte zwischen B und J als einem Endpunkte und mit dem anderen Endpunkte in C , von der Geraden $Y = 0$ enthält es die Strecke A_0OA . Danach enthält es keinen Gitterpunkt außer O und A, A_0, C, C_0 . Da ferner B außerhalb dieses Parallelogramms liegt und $AC = OB$ ist, so ist AC größer als die Hälfte der A und C enthaltenden Seite von $\mathfrak{P}(\varrho, \sigma)$ und befindet sich daher C näher an $S (= N)$ und A weiter von S als die Mitte $\xi = \frac{\varepsilon\varrho}{2}$, $\eta = \frac{\sigma}{2}$ dieser Seite; man hat also $\lambda' < \frac{\varrho}{2} < \lambda$, $\mu' > \frac{\sigma}{2} > \mu$.

Jetzt sei $OM = ON$, so daß M mit N zusammenfällt. Nehmen wir zudem $\mu > 0$ an, so daß A nicht Ecke in $\mathfrak{P}(\varrho, \sigma)$ wird, so ist $GO > A_0O$ und daher wieder $BJ > JC$, $\lambda - \lambda' > \lambda'$, und gelten alle Überlegungen wie vorhin, nur daß außer A, A_0, C, C_0 noch die Gitterpunkte B und B_0 auf die Begrenzung von $\mathfrak{P}(\varrho, \sigma)$ zu liegen kommen. Weil OB parallel AC ist, wird dabei B Mitte einer Seite von $\mathfrak{P}(\varrho, \sigma)$, also $\lambda - \lambda' = \frac{\varrho}{2}$, $\mu' - \mu = \frac{\sigma}{2}$, dagegen ist, weil $OB = AC$ und weder A noch C eine Ecke von $\mathfrak{P}(\varrho, \sigma)$ ist, weder C noch A Mitte einer Seite von $\mathfrak{P}(\varrho, \sigma)$; man hat wieder $\lambda' < \frac{\varrho}{2} < \lambda$, $\mu' > \frac{\sigma}{2} > \mu$.

Hat man schließlich $OM = ON$ und dazu $\mu = 0$, so ist A_0OA Diagonale von $\mathfrak{P}(\varrho, \sigma)$ und fällt G mit A_0 zusammen. Dann folgt $BJ = JC = \frac{1}{2}OA$, also $\lambda - \lambda' = \lambda'$ und weiter $JN = OJ$, $CN = AC = OB$. Unter diesen Umständen reicht $\mathfrak{P}(\varrho, \sigma)$ mit der Ecke S gerade an $Y = 2$ heran, es enthält wieder im Inneren außer O keinen Gitterpunkt, aber nicht bloß B , sondern auch C ist darin Mitte einer Seite. Für B wie für C gilt dann $|\xi\eta| = \frac{1}{2}$. Es wäre dieses derjenige Fall, wo $\xi\eta$ sich als äquivalent mit der Form $\frac{1}{2}(X^2 - Y^2)$ erweist, ein Fall, der vorweg ausgeschlossen wurde.

Aus den Betrachtungen in § 1 folgt nunmehr in jedem Falle, daß der Gitterpunkt A' den Forderungen des Satzes II entspricht. Zur Bestimmung dieses Gitterpunktes $x = p'$, $y = q'$ aus dem Gitterpunkte A hat sich sogleich die folgende Regel herausgestellt:

Man bezeichne mit g die größte in $-\frac{\lambda}{\lambda}$ enthaltene ganze Zahl und setze $h = g$ oder $h = g + 1$, je nachdem

$$(-\bar{\lambda} - \lambda g)(\bar{\mu} + \mu g) < \text{oder} \geq \frac{1}{2}$$

ist. Dann hat man $p' = r + ph$, $q' = s + qh$.

2. Verändern wir die Parameter ϱ, σ des in 1. betrachteten Parallelogramms $\mathfrak{P}(\varrho, \sigma)$ in der Weise, daß σ verkleinert wird, also S und S_0 näher an O heranrücken, daß aber die Seiten noch fortwährend durch A, A_0 (und F, F_0) gehen, so erhalten wir, so lange die Abnahme von σ eine gewisse Grenze nicht erreicht, ein neues Parallelogramm mit $\xi = 0, \eta = 0$ als Diagonalen, welches offenbar keine anderen Gitterpunkte außer A_0, O, A enthält. Daraus geht der Satz hervor:

Satz III. Ist ein Gitterpunkt $x = p, y = q$ so beschaffen, daß die Zahlen p, q relativ prim sind und dafür $|\xi\eta| < \frac{1}{2}$ ausfällt, so kann man stets Parallelogramme

$$\left| \frac{\xi}{\varrho} \right| + \left| \frac{\eta}{\sigma} \right| \leq 1$$

konstruieren, welche den Gitterpunkt auf der Begrenzung liegen haben und außer den drei Punkten $x, y = p, q; 0, 0; -p, -q$ überhaupt keinen Gitterpunkt enthalten.

Wenn andererseits für einen Gitterpunkt $x = p, y = q$ ein Parallelogramm der hier bezeichneten Art existiert, so zeigt der Beweis zum Satze I, daß umgekehrt stets p, q relativ prim sind und dafür $|\xi\eta| < \frac{1}{2}$ ausfällt.

Auf Grund dieses Satzes III beweisen wir über das Verhältnis der in 1. mit A und A' bezeichneten zwei Gitterpunkte den folgenden wichtigen

Zusatz: Es kann keinen von A, A_0, A', A'_0 verschiedenen Gitterpunkt x, y geben, für den ebenfalls x, y relativ prim sind und ebenfalls $|\xi\eta| < \frac{1}{2}$, dabei aber $\lambda \geq |\xi| \geq \lambda'$ wäre.

Beweis. Nehmen wir an, es existierte ein Gitterpunkt A^* der hier bezeichneten Art, und es sei für ihn $|\xi| = \lambda^*, |\eta| = \mu^*$. Wir bezeichnen (Fig. 4) mit E den Punkt $\xi = \lambda, \eta = \mu$, mit E' den Punkt $\xi = \lambda', \eta = \mu'$, mit R und S die Schnittpunkte der Geraden EE' mit $\eta = 0$ und $\xi = 0$, weiter mit L den Punkt $\xi = \sigma, \eta = 0$, mit L' den Punkt $\xi = \lambda', \eta = 0$, endlich mit E^* den Punkt $\xi = \lambda^*, \eta = \mu^*$. Nach dem Satze III müßte es zufolge unserer Voraussetzungen möglich sein, ein Parallelogramm $\mathfrak{P}(\varrho^*, \sigma^*)$ mit $\xi = 0, \eta = 0$ als Diagonalen zu konstruieren, dessen Begrenzung den Punkt A^* aufnimmt, welches aber weder A noch A' enthielte. Sind R^*, S^* die Ecken $\xi = \varrho^*, \eta = 0$ und $\xi = 0, \eta = \sigma^*$ dieses Parallelogramms, so enthält die Strecke R^*S^* den Punkt E^* , es darf aber weder E , noch E' dem Dreiecke OR^*S^* an-

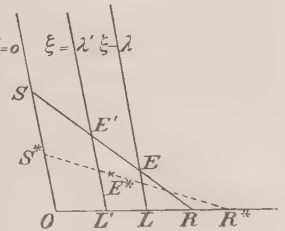


Fig. 4.

gehören. Da nun nach Voraussetzung E^* in dem Parallelstreifen $\lambda' \leq \xi \leq \lambda$ liegt und noch dafür $\eta \geq 0$ ist, müßte danach E^* sich notwendig im Viereck $L'LEE'$ und dabei nicht auf der Seite EE' desselben befinden. Aber das Parallelogramm RSR_0S_0 enthält im Inneren O als einzigen Gitterpunkt, danach kann E^* nicht in $L'LEE'$ bei Ausschluß der Strecke EE' liegen. Ein Gitterpunkt, wie wir ihn in A^* angenommen haben, kann also nicht existieren.

In entsprechender Weise leuchtet ein, daß es keinen von A, A_0, A', A'_0 verschiedenen Gitterpunkt geben kann, für den ebenfalls x, y relativ prim sind und ebenfalls $|\xi\eta| < \frac{1}{2}$, dabei aber $\mu \leq |\eta| \leq \mu'$ wäre.

3. Durch die aus A und A' entnommene Substitution

$$T) \quad x = pX + p'Y, \quad y = qX + q'Y,$$

deren Determinante $pq' - qp' = \pm 1$ ist, geht die Form $\xi\eta$ in eine äquivalente Form

$$\varphi X^2 + \chi XY + \psi Y^2 = (\varepsilon\lambda X + \varepsilon'\lambda'Y)(\mu X + \mu'Y)$$

über.

Man erhält sodann zunächst die Gleichung $\chi^2 - 4\varphi\psi = 1$.

Sodann ist $\varphi = \varepsilon\lambda\mu$, $\psi = \varepsilon'\lambda'\mu'$ und gilt $|\varphi| < \frac{1}{2}$, $|\psi| < \frac{1}{2}$.

Weiter hat man, wenn $\varepsilon' = -\varepsilon$ ist, $\lambda\mu' + \mu\lambda' = 1$, $\lambda > \lambda' \geq 0$, $\mu' > \mu \geq 0$ und $\chi = \varepsilon(\lambda\mu' - \mu\lambda') = \varepsilon(1 - 2\mu\lambda')$; also ist in diesem Falle $0 < \varepsilon\chi \leq 1$. Wenn dagegen $\varepsilon' = \varepsilon$ ist, hat man $\lambda\mu' - \mu\lambda' = 1$, $\lambda > 2\lambda' > 0$, $\mu' > 2\mu \geq 0$ und $\chi = \varepsilon(\lambda\mu' + \mu\lambda') = \varepsilon(1 + 2\mu\lambda')$; da hier $2\mu\lambda' < \mu\lambda' < \frac{1}{2}$ ist, folgt also $1 \leq \varepsilon\chi < \frac{3}{2}$. Die Relation $(\lambda - \lambda')(\mu' - \mu) \geq \frac{1}{2}$, die in diesem zweiten Falle besteht, ergibt $\varepsilon(\chi - \varphi - \psi) \geq \frac{1}{2}$. — Die Vorzeichen ε und ε' entnimmt man hiernach bereits aus dem Werte von χ allein, außer wenn gerade $\chi = \pm 1$ (also $\varphi = 0$ oder $\psi = 0$) ist.

Da aus $pq' - qp' = \pm 1$ schon hervorgeht, daß p und q einerseits und andererseits p' und q' relativ prim sind, so sind die besonderen Umstände, welche hier für die zwei Gitterpunkte $A(\xi = \varepsilon\lambda, \eta = \mu)$ und $A'(\xi = \varepsilon'\lambda', \eta = \mu')$ gelten, völlig zusammengefaßt in den Beziehungen: $\lambda > 0$; $\mu > 0$, $\varepsilon = \pm 1$ oder $\mu = 0$, $\varepsilon = 1$; $pq' - qp' = \varepsilon$, ferner *entweder*:

$$\varepsilon' = -\varepsilon, \quad \lambda > \lambda' \geq 0, \quad \lambda\mu < \frac{1}{2}, \quad \lambda'\mu' < \frac{1}{2}$$

oder:

$$\varepsilon' = \varepsilon, \quad \lambda > \lambda' > 0, \quad \lambda\mu < \frac{1}{2}, \quad (\lambda - \lambda')(\mu' - \mu) \geq \frac{1}{2}.$$

Bemerkenswert ist noch, daß bei dieser zweiten Reihe von Bedingungen für $\varepsilon' = \varepsilon$ die Ungleichung $\lambda\mu < \frac{1}{2}$ eine Folge der übrigen Ungleichungen

ist. Denn man hat hier $\lambda\mu' - \lambda'\mu = 1$. Aus

$$(\lambda - \lambda')(\mu' - \mu) \geq \frac{1}{2}$$

und $\lambda > \lambda'$ folgt zuvörderst $\mu' > \mu$; sodann kann diese Ungleichung ersetzt werden durch

$$\lambda\mu' + \lambda'\mu - \lambda\mu - \lambda'\mu' \geq \frac{1}{2}(\lambda\mu' - \lambda'\mu),$$

d. i.

$$\frac{1}{2}\lambda\mu' + \frac{3}{2}\lambda'\mu - \lambda\mu - \lambda'\mu' \geq 0$$

oder

$$\frac{1}{2}(\lambda - \lambda')(\mu' - 2\mu) \geq \frac{1}{2}\lambda'(\mu' - \mu).$$

Daraus folgt $\mu' - 2\mu > 0$. Ersetzt man weiter hier $\lambda\mu'$ durch $1 + \lambda'\mu$, so entsteht endlich

$$\frac{1}{2} + 2\lambda'\mu \geq \lambda\mu + \lambda'\mu', \quad \text{also} \quad \frac{1}{2} \geq \lambda\mu + \lambda'(\mu' - 2\mu),$$

und daraus entnimmt man in der Tat $\frac{1}{2} > \lambda\mu$.

§ 3.

Fassen wir die Resultate des § 1 und § 2 zusammen und nehmen die Sätze hinzu, welche daraus bei Vertauschung der Rollen von ξ und η , bei Ersetzung von ξ, η durch $\eta, -\xi$, hervorgehen, so entsteht der folgende

Satz IV. *Es seien $\xi = \alpha x + \beta y$, $\eta = \gamma x + \delta y$ zwei lineare Formen mit beliebigen reellen Koeffizienten und einer Determinante $\alpha\delta - \beta\gamma = 1$; jedoch sei die Form $\xi\eta$ in x, y nicht äquivalent mit der Form XY oder der Form $\frac{1}{2}(X^2 - Y^2)$ in X, Y . Alsdann lassen sich die sämtlichen Systeme von ganzen Zahlen x, y , für welche x, y relativ prim sind und $|\xi\eta| < \frac{1}{2}$ und zudem $\eta > 0$, bzw. $\eta = 0$, $\xi > 0$ ist, in eine Reihe nach wachsendem Werte η ordnen. Dabei sind sie zugleich nach abnehmendem Werte $|\xi|$ geordnet.*

Für je zwei aufeinanderfolgende Systeme $x = p, y = q$ und $x = p', y = q'$ in der Reihe gilt dann stets

$$pq' = qp' = \pm 1.$$

Diese Reihe weist ein bestimmtes erstes System auf, wofür $\eta = 0, \xi > 0$, also $\frac{x}{\delta} = \frac{y}{-\gamma}$ und > 0 ist, wenn $\frac{\delta}{-\gamma}$ rational ist; sie weist ein bestimmtes letztes System auf, wofür $\xi = 0, \eta > 0$, also $\frac{x}{-\beta} = \frac{y}{\alpha}$ und > 0 ist, wenn $\frac{-\beta}{\alpha}$ rational ist. Sie ist ohne ein erstes System, wenn $\frac{\delta}{-\gamma}$ irrational ist,

ohne ein letztes System, wenn $\frac{-\beta}{\alpha}$ irrational ist; sie ist nach Anfang und Ende hin unbegrenzt, wenn sowohl $\frac{\delta}{\gamma}$ wie $\frac{-\beta}{\alpha}$ irrational sind.

Ist die Reihe ohne ein letztes System, so konvergiert in ihrem Verlaufe $|\xi|$ nach Null und wächst η über jede Grenze; ist sie ohne ein erstes System, so wächst bei umgekehrter Folge der Systeme in ihr $|\xi|$ über jede Grenze und konvergiert η nach Null.

Diese zuletzt erwähnte Tatsache folgt aus dem Umstande, daß überhaupt nur für eine endliche Anzahl von Systemen der Reihe $|\xi|$ oder η zwischen gegebenen positiven Grenzen liegen kann. Denn soll etwa $\varrho_1 \geq |\xi| \geq \varrho_0 > 0$ sein, so folgt aus $|\xi\eta| < \frac{1}{2}$ und $|\xi| \geq \varrho_0$ noch $|\eta| < \frac{1}{2\varrho_0}$; in einem Parallelogramme $|\xi| \leq \varrho_1$, $|\eta| \leq \frac{1}{2\varrho_0}$ liegen aber stets nur eine endliche Anzahl von Gitterpunkten x, y .

Die Reihe der hier in Betracht kommenden Gitterpunkte x, y , nach wachsendem Werte ihres η geordnet, soll die *Kette* zu den Formen ξ, η heißen, ein einzelner Gitterpunkt $x = p, y = q$ daraus ein *Kettenglied*, ferner die mittels zweier aufeinanderfolgender Kettenglieder $x = p, y = q$ und $x = p', y = q'$ gebildete Substitution

$$x = pX + p'Y, \quad y = qX + q'Y$$

eine *Substitution der Kette* heißen.

Wir bezeichnen die nacheinander auftretenden Kettenglieder mit

$$p_i, q_i \quad (i = \dots - 2, -1, 0, 1, 2, \dots),$$

wobei wir, wenn ein erstes Glied vorhanden ist, diesem Gliede und anderenfalls einem beliebig gewählten Gliede den Index 0 erteilen wollen. Für $x = p_i, y = q_i$ setzen wir $\xi = \varepsilon_i \lambda_i, \eta = \mu_i$, so daß $\mu_i \geq 0, \lambda_i \geq 0, \varepsilon_i = \pm 1$ sei. Ferner schreiben wir allgemein, soweit die Indizes in Betracht kommen,

$$\frac{\varepsilon_i}{\varepsilon_{i-1}} = \vartheta_i.$$

Für ein etwa vorhandenes erstes Glied ist $\varepsilon_0 = 1$. Für einen etwa vorhandenen letzten Index $i = w$, wobei dann $\lambda_w = 0$ wäre, denken wir uns $\vartheta_w = -1$ gewählt. Die Substitution

$$x = p_{i-1}X_i + p_iY_i, \quad y = q_{i-1}X_i + q_iY_i$$

oder kurz $\begin{pmatrix} p_{i-1} & p_i \\ q_{i-1} & q_i \end{pmatrix}$ bezeichnen wir mit T_i .

Man hat dann allgemein:

$$\lambda_{i-1} > \lambda_i, \quad \mu_{i-1} < \mu_i,$$

ferner

$$(1) \quad \lambda_{i-1}\mu_i - \vartheta_i\mu_{i-1}\lambda_i = 1, \quad p_{i-1}q_i - q_{i-1}p_i = \varepsilon_{i-1}.$$

Die am Schlusse von § 2, 1. entwickelte Regel, welche überhaupt dazu verhilft, aus einem Kettengliede das nächstfolgende abzuleiten, ergibt einen einfachen Zusammenhang zwischen drei aufeinanderfolgenden Kettengliedern: p_{i-1} , q_{i-1} ; p_i , q_i ; p_{i+1} , q_{i+1} . Wir können p_i , q_i mit dem Gitterpunkte p , q (ε_i mit ε , λ_i mit λ) und p_{i+1} , q_{i+1} mit p' , q' aus § 2 identifizieren. Für die Zahlen r , s dort, welche der Bedingung $p_i s - q_i r = \varepsilon_i$ zu genügen haben, kann man dann mit Rücksicht auf (1): $s = -\vartheta_i q_{i-1}$, $r = -\vartheta_i p_{i-1}$ einführen, und dabei wird

$$\xi = \varepsilon_i \bar{\lambda} = -\vartheta_i \varepsilon_{i-1} \lambda_{i-1}.$$

Also ist dann $\bar{\lambda}$ durch $-\lambda_{i-1}$ zu ersetzen. Dadurch entsteht die folgende Regel:

Man bezeichne mit g_i die größte in $\frac{\lambda_{i-1}}{\lambda_i}$ enthaltene ganze Zahl und setze $h_i = g_i$ oder $= g_i + 1$, je nachdem

$$(\lambda_{i-1} - g_i \lambda_i) (g_i \mu_i - \vartheta_i \mu_{i-1}) < \text{oder} \geq \frac{1}{2}$$

ist, dann gilt

$$p_{i+1} = -\vartheta_i p_{i-1} + h_i p_i, \quad q_{i+1} = -\vartheta_i q_{i-1} + h_i q_i$$

und überdies wird $\vartheta_{i+1} = -1$ im ersten, $= 1$ im zweiten Falle.

Man erhält hiernach

$$(2) \quad \begin{pmatrix} p_{i-1} & p_i \\ q_{i-1} & q_i \end{pmatrix} \begin{pmatrix} 0 & -\vartheta_i \\ 1 & h_i \end{pmatrix} = \begin{pmatrix} p_i & p_{i+1} \\ q_i & q_{i+1} \end{pmatrix},$$

$$\begin{pmatrix} \varepsilon_{i-1} \lambda_{i-1} & \varepsilon_i \lambda_i \\ \mu_{i-1} & \mu_i \end{pmatrix} \begin{pmatrix} 0 & -\vartheta_i \\ 1 & h_i \end{pmatrix} = \begin{pmatrix} \varepsilon_i \lambda_i & \varepsilon_{i+1} \lambda_{i+1} \\ \mu_i & \mu_{i+1} \end{pmatrix}.$$

Es wird also

$$p_{i-1} X_i + p_i Y_i = p_i X_{i+1} + p_{i+1} Y_{i+1},$$

$$q_{i-1} X_i + q_i Y_i = q_i X_{i+1} + q_{i+1} Y_{i+1}$$

vermöge

$$X_i = -\vartheta_i Y_{i+1}, \quad Y_i = X_{i+1} + h_i Y_{i+1}.$$

Setzt man allgemein $-\frac{X_i}{Y_i} = t_i$, so gilt daher weiter

$$(3) \quad \frac{-p_{i-1} t_i + p_i}{-q_{i-1} t_i + q_i} = \frac{-p_i t_{i+1} + p_{i+1}}{-q_i t_{i+1} + q_{i+1}},$$

während

$$(4) \quad t_i = \frac{\vartheta_i}{h_i - t_{i+1}}$$

ist. Dabei ist zu bemerken, daß Zähler und Nenner der rechten Seite von (3) genau die Ausdrücke sind, die bei der naturgemäßen Umformung des Zählers oder des Nenners der linken Seite je in einen Quotienten zweier ganzer Funktionen von t_{i+1} als die Zähler dieser beiden Quotienten erscheinen. Aus (3) und (4) erhält man sogleich allgemeiner:

$$(5) \quad \frac{-p_{i-1}t_i + p_i}{-q_{i-1}t_i + q_i} = \frac{-p_{k-1}t_k + p_k}{-q_{k-1}t_k + q_k}, \quad i < k,$$

wenn

$$(6) \quad t_i = \frac{\vartheta_i}{h_i - \frac{\vartheta_{i+1}}{h_{i+1} - \frac{\vartheta_{i+2}}{\ddots - \frac{\vartheta_{k-1}}{h_{k-1} - t_k}}}}$$

ist, und dabei gilt über die Entstehung von Zähler und Nenner der rechten Seite in (5) aus Zähler und Nenner der linken Seite eine ganz entsprechende Bemerkung wie bei den Beziehungen (3) und (4).

Ebenso wie ξ, η haben $\eta, -\xi$ die Determinante 1. Die Kette zu $\eta, -\xi$ ist im wesentlichen die umgekehrte Kette zu ξ, η . Durchläuft p_i, q_i die Gitterpunkte der Kette zu ξ, η in umgekehrter Reihenfolge, d. h. so daß die Indizes i abnehmen, so hat man dabei in den Systemen $x = -\varepsilon_i p_i, y = -\varepsilon_i q_i$, (für welche $-\xi \geq 0$ ausfällt), die Glieder der Kette zu $\eta, -\xi$. Über den Fortgang in dieser letzteren Kette sei noch die aus (2) folgende Beziehung:

$$(7) \quad \begin{pmatrix} -\varepsilon_{i+1}p_{i+1}, & -\varepsilon_i p_i \\ -\varepsilon_{i+1}q_{i+1}, & -\varepsilon_i q_i \end{pmatrix} \begin{pmatrix} 0, & -\vartheta_{i+1} \\ 1, & h_i \end{pmatrix} = \begin{pmatrix} -\varepsilon_i p_i, & -\varepsilon_{i-1}p_{i-1} \\ -\varepsilon_i q_i, & -\varepsilon_{i-1}q_{i-1} \end{pmatrix}$$

angemerkt.

§ 4.

An die Sätze in § 2 schließt sich folgendes Theorem an:

Satz V. Sind $\xi = \alpha x + \beta y, \eta = \gamma x + \delta y$ zwei lineare Formen mit beliebigen reellen Koeffizienten und einer Determinante $\alpha\delta - \beta\gamma = 1$ und sind ξ_0, η_0 irgendwelche gegebene reelle Größen, so gibt es stets ganze Zahlen x, y , für welche

$$|(\xi - \xi_0)(\eta - \eta_0)| \leq \frac{1}{4}$$

ausfällt.

Beweis. Betrachten wir vorweg die Fälle, daß es eine ganzzahlige Substitution mit einer Determinante ± 1 gibt, durch welche die Form $\xi\eta$ der Variablen x, y in die Form XY oder in die Form $\frac{1}{2}(X^2 - Y^2)$ der neuen Variablen X, Y übergehe. Dem Wertsysteme $\xi = \xi_0, \eta = \eta_0$ entspreche dabei das System $X = X_0, Y = Y_0$, so erweist sich vermöge jener Substitution

$$(\xi - \xi_0)(\eta - \eta_0) = (X - X_0)(Y - Y_0), \text{ bzw. } = \frac{1}{2}((X - X_0)^2 - (Y - Y_0)^2).$$

Bestimmen wir nun X und Y als ganze Zahlen so, daß $|X - X_0| \leq \frac{1}{2}, |Y - Y_0| \leq \frac{1}{2}$ ist, so wird im ersten Falle, wo $\xi\eta$ äquivalent XY ist,

$|(\xi - \xi_0)(\eta - \eta_0)| \leq \frac{1}{4}$. Dabei ist zu bemerken, daß das Gleichheitszeichen hier unter gewissen Umständen wirklich in Betracht kommt, nämlich wenn sowohl X_0 wie Y_0 gleich ganzen Zahlen vermehrt um $\frac{1}{2}$ sind. — Im zweiten Falle, wo $\xi\eta$ äquivalent $\frac{1}{2}(X^2 - Y^2)$ ist, stellt sich sogar $|(\xi - \xi_0)(\eta - \eta_0)| \leq \frac{1}{8}$ heraus.

Nunmehr schließen wir die Fälle aus, daß $\xi\eta$ äquivalent mit XY oder mit $\frac{1}{2}(X^2 - Y^2)$ ist.

Betrachten wir irgendeine Substitution

$$x = pX + p'Y, \quad y = qX + q'Y$$

der zu ξ, η gehörigen Kette. Wir bezeichnen den Gitterpunkt p, q mit A , den Gitterpunkt p', q' mit A' . Nach § 2 haben wir ein ganz bestimmtes Parallelogramm RSR_0S_0 oder $\mathfrak{P}(\rho, \sigma)$: $\left|\frac{\xi}{\rho}\right| + \left|\frac{\eta}{\sigma}\right| \leq 1$, dessen Berandung sowohl A , wie A' aufnimmt. Der größeren Anschaulichkeit wegen wollen wir in der Zeichnungsebene die Koordinaten x, y derart interpretieren, daß dieses Parallelogramm $\mathfrak{P}(\rho, \sigma)$ ein *Quadrat* im gewöhnlichen Sinne wird. Für p, q sei $\xi = \varepsilon\lambda, \eta = \mu, (\lambda \geq 0, \mu \geq 0, \varepsilon = \pm 1)$, für p', q' sei $\xi = \varepsilon'\lambda', \eta = \mu', (\lambda' \geq 0, \mu' \geq 0, \varepsilon' = \pm 1)$. Wir haben nun die beiden, auch in § 2 unterschiedenen Fälle $\varepsilon' = -\varepsilon$ und $\varepsilon' = \varepsilon$ gesondert zu untersuchen.

1°. Es sei zunächst $\varepsilon' = -\varepsilon$ und $\lambda > \lambda' \geq 0, \mu' > \mu \geq 0$. Ein gleichzeitiges Eintreten von $\mu = 0$ und $\lambda' = 0$ ist dadurch ausgeschlossen, daß $\xi\eta$ nicht äquivalent XY sein soll. Es sei M die Seitenmitte $\xi = \frac{\varepsilon\rho}{2}, \eta = \frac{\sigma}{2}$ und M' die Seitenmitte $\xi = -\frac{\varepsilon\rho}{2}, \eta = \frac{\sigma}{2}$ in $\mathfrak{P}(\rho, \sigma)$. Nach den Bemerkungen in § 2 kommen A und A' derart auf der Berandung von $\mathfrak{P}(\rho, \sigma)$ zu liegen, daß bei einer Umlaufung derselben in einem gewissen Sinne sich $A, M; (S), A', M', A_0, M_0, (S_0), A'_0, M'_0$ folgen (s. Fig. 5; um die Figur nicht zu komplizieren, sind darin die Seiten von $\mathfrak{P}(\rho, \sigma)$ nicht ausgezogen); A ist von M , A' von M' verschieden.

Da wir die Rollen von ξ und η vertauschen, anstatt ξ, η auch $\eta, -\xi$ zugrunde legen können, so dürfen wir noch die Annahme $\lambda'\mu' \geq \lambda\mu$ machen; (weil nicht zugleich $\lambda' = 0, \mu = 0$ sein kann, wird dann gewiß $\lambda' > 0$, also A' von S verschieden sein). Da auf jeder Seite des Quadrates $\mathfrak{P}(\rho, \sigma)$ der Wert $\left|\frac{\xi\eta}{\rho\sigma}\right|$ stets von der Mitte der Seite nach ihren Enden hin abnimmt und dabei auf allen vier Seiten gleichen Wert erhält bei der nämlichen *Entfernung* von der Seitenmitte, den Begriff der Entfernung

im gewöhnlichen Sinne genommen, so läuft jene Annahme darauf hinaus, daß $A'M' \leq AM$ sein soll.

Wir konstruieren weiter das Quadrat $\mathfrak{P}\left(\frac{\rho}{2}, \frac{\sigma}{2}\right)$, dessen Ecken auf $\eta = 0$, $\xi = 0$ halb so große Entfernungen von O haben wie R, R_0, S, S_0 . Die Berandung von $\mathfrak{P}\left(\frac{\rho}{2}, \frac{\sigma}{2}\right)$ nimmt die Punkte $X = \pm \frac{1}{2}$, $Y = 0$ und $X = 0$, $Y = \pm \frac{1}{2}$ auf. Legen wir sodann um jeden einzigen Gitterpunkt als Mittelpunkt ein Quadrat, welches dem Quadrate $\mathfrak{P}\left(\frac{\rho}{2}, \frac{\sigma}{2}\right)$ gleich und in den Seiten parallel gestellt ist, so ergeben diese Quadrate das Bild der in Fig. 5 mit ausgezogener Umrandung gezeichneten Quadrate. Die einzelnen Gitterpunkte sind in dieser Figur durch ihre Koordinaten X, Y angedeutet. Das erste Quadrat $\mathfrak{P}\left(\frac{\rho}{2}, \frac{\sigma}{2}\right)$ um den Nullpunkt stößt mit Stücken seiner Seiten an die Quadrate mit den Mittelpunkten A, A', A_0, A'_0 , während es die übrigen Quadrate überhaupt nicht trifft. Danach liegen jene Quadrate offenbar so, daß keine zwei ineinander eingreifen, und zwischen sich lassen sie noch lauter gleiche und parallel liegende Lücken in Form von Rechtecken, welche die einzelnen Punkte $X + \frac{1}{2}$, $Y + \frac{1}{2}$ mit ganzzahligen X, Y zu Mittelpunkten haben.

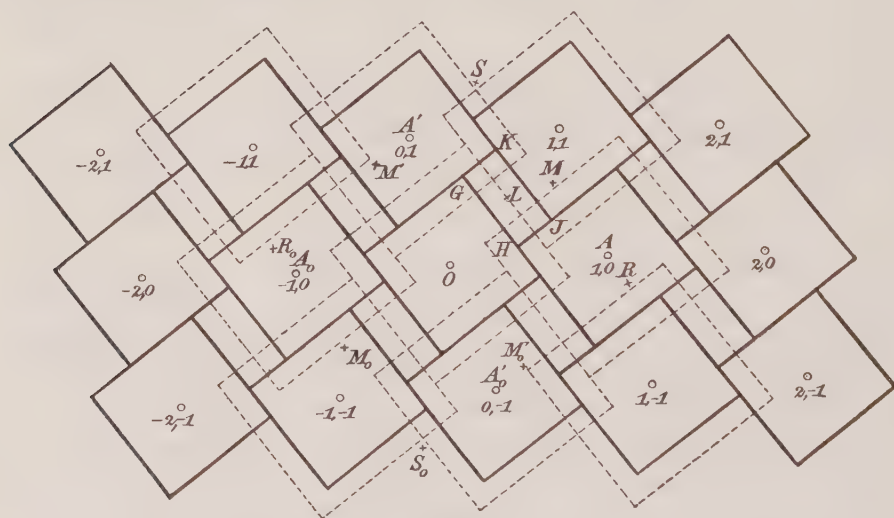


Fig. 5.

Es sei beispielsweise $GHJK$ die rechteckige Lücke mit dem Mittelpunkte $X = \frac{1}{2}$, $Y = \frac{1}{2}$ oder L , so daß die Seiten GH, HJ, JK, KG sich an die Quadrate mit den Mittelpunkten $X, Y = 0, 0; 1, 0; 1, 1; 0, 1$

anlegen. Da AH , $M'GM$, $A'K$ sämtlich parallel der Linie $\eta = 0$ sind, so ist $GH = MA \leq MS$ und $GK = M'A' < M'S$, also $GH \geq GK$ und überdies $\frac{1}{2} RS \geq GH$, $\frac{1}{2} RS > GK$; (man beachte, daß A' nicht in S fällt). In jeder der rechteckigen Lücken ist daher eine Seite kleiner, die andere höchstens so groß als die Seite des Quadrates $\mathfrak{P}\left(\frac{\varrho}{2}, \frac{\sigma}{2}\right)$.

Unter dem *Inhalt* einer Figur wollen wir den Wert des Integrales $\iint dX dY$ über ihre Fläche verstehen. Der Inhalt des Quadrates $\mathfrak{P}\left(\frac{\varrho}{2}, \frac{\sigma}{2}\right)$ ist dann, weil $\alpha\delta - \beta\gamma = 1$, $pq' - qp' = \pm 1$ ist, gleich $\frac{1}{2} \varrho\sigma$ und der Inhalt des Rechteckes $GHJK$ ist $< \frac{1}{2} \varrho\sigma$. Nun kommt in der ganzen Ebene auf jeden Gitterpunkt $X = X^*$, $Y = Y^*$ einerseits ein Parallelogramm $-\frac{1}{2} \leq X - X^* \leq \frac{1}{2}$, $-\frac{1}{2} \leq Y - Y^* \leq \frac{1}{2}$ vom Inhalte 1, wobei diese Parallelogramme die ganze Ebene lückenlos erfüllen, ohne gegenseitig ineinander einzudringen, andererseits kommt hier auf jeden Gitterpunkt X^* , Y^* ein Quadrat mit dem Mittelpunkte X^* , Y^* vom Inhalte $\frac{1}{2} \varrho\sigma$ und ein Rechteck mit dem Mittelpunkte $X^* + \frac{1}{2}$, $Y^* + \frac{1}{2}$ von einem gewissen Inhalte $< \frac{1}{2} \varrho\sigma$, und alle diese Quadrate und Rechtecke erfüllen ebenfalls die ganze Ebene lückenlos, ohne gegenseitig ineinander einzudringen. Danach folgt offenbar

$$(1) \quad \frac{1}{2} \varrho\sigma < 1 < 2 \cdot \frac{1}{2} \varrho\sigma.$$

Es verhalte sich die Entfernung des Punktes O von der Geraden GH zu der des Punktes L von dieser Geraden, also $\frac{1}{2} M'S : \frac{1}{2} M'A'$ wie $1 : \kappa - 1$, dabei ist $1 < \kappa < 2$. Konstruieren wir dann das Quadrat $\mathfrak{P}\left(\frac{\kappa\varrho}{2}, \frac{\kappa\sigma}{2}\right)$ mit O als Mittelpunkt, dessen Seiten denen von $\mathfrak{P}\left(\frac{\varrho}{2}, \frac{\sigma}{2}\right)$ parallel sind, so geht dessen Berandung durch L und wird deshalb dieses Quadrat eine Hälfte der Lücke $GHJK$ vollständig überdecken. Legen wir nun um jeden Gitterpunkt als Mittelpunkt ein diesem Quadrate $\mathfrak{P}\left(\frac{\kappa\varrho}{2}, \frac{\kappa\sigma}{2}\right)$ gleiches und parallel gestelltes Quadrat, so werden daher diese Quadrate zweiter Art, die in Fig. 5 mit gestrichelter Berandung gezeichnet sind, jedenfalls die ganze Ebene, ohne Lücken zu lassen, überdecken. Also wird der Punkt, für den die Bestimmungsstücke ξ, η gleich den gegebenen Werten ξ_0, η_0 sind, in wenigstens eines dieser Quadrate fallen müssen; für den Gitterpunkt x, y , welcher der Mittelpunkt des betreffenden Quadrates ist, gilt dann

$$(2) \quad \left| \frac{\xi - \xi_0}{\varrho} \right| + \left| \frac{\eta - \eta_0}{\sigma} \right| \leq \frac{\kappa}{2}.$$

In das Innere des Quadrates $\mathfrak{P}\left(\frac{\kappa\rho}{2}, \frac{\kappa\sigma}{2}\right)$ dringen von allen übrigen dieser Quadrate zweiter Art nur die vier Quadrate mit den Mittelpunkten A, A_0, A', A_0' . Dabei liegen diese vier Quadrate selbst völlig auseinander, auch reichen sie nicht bis an den Nullpunkt O heran. Danach zeigt sich, daß die Gesamtheit dieser Quadrate zweiter Art die Ebene so überdecken, daß sie kein Gebiet mehr als *zweifach* überlagern, während noch gewisse $\sqcup$ förmige Partien mit den einzelnen Gitterpunkten als Mittelpunkten vorhanden sind, die jedesmal nur je einem dieser Quadrate angehören. Infolgedessen ist der Inhalt des Quadrates $\mathfrak{P}\left(\frac{\kappa\rho}{2}, \frac{\kappa\sigma}{2}\right)$ notwendig < 2 , d. h. man hat

$$(3) \quad \frac{1}{2} \kappa^2 \rho \sigma < 2.$$

Aus (2) folgt, weil das geometrische Mittel zweier Beträge nicht größer als ihr arithmetisches Mittel ist,

$$\left| \frac{(\xi - \xi_0)(\eta - \eta_0)}{\rho \sigma} \right| \leq \left(\frac{\kappa}{4} \right)^2,$$

und daraus mit Rücksicht auf (3)

$$|(\xi - \xi_0)(\eta - \eta_0)| < \frac{1}{4}.$$

2°. Zweitens sei $\varepsilon' = \varepsilon$ und $\lambda > \lambda' > 0$, $\mu' > \mu \geq 0$. Die Punkte A und A' werden hier von einer und derselben Seite von $\mathfrak{P}(\rho, \sigma)$ aufgenommen, wobei weder A noch A' in die Mitte der Seite fallen und A , aber nicht A' eine Ecke sein kann. Die Länge der betreffenden Seite ist $\leq 2AA'$.

Gehen wir nunmehr zu dem Quadrat $\mathfrak{P}\left(\frac{\rho}{2}, \frac{\sigma}{2}\right)$ über und konstruieren um jeden Gitterpunkt als Mittelpunkt ein diesem gleiches und parallel gestelltes Quadrat, so liefern diese Quadrate das Bild der in Fig. 6 mit ausgezogener Umrandung gezeichneten Quadrate. Die Gitterpunkte sind dort durch ihre Koordinaten X, Y bezeichnet. Diese verschiedenen Quadrate nun greifen nicht ineinander ein und lassen zwischen sich im allgemeinen wieder lauter gleiche und parallel gelagerte rechteckige Lücken mit den einzelnen Punkten $X + \frac{1}{2}, Y + \frac{1}{2}$ für ganzzahlige X, Y als Mittelpunkten. Diese Lücken kommen nur zum Fortfall, wenn auch der Gitterpunkt $X = -1, Y = 1$ auf die Begrenzung von $\mathfrak{P}(\rho, \sigma)$ fällt, (wenn $(\lambda - \lambda')(\mu' - \mu) = \frac{1}{2}$ ist). Es sei, wenn nicht dieser Spezialfall statthat, $G H J K$ die rechteckige Lücke mit dem Mittelpunkte L : $X = \frac{1}{2}, Y = \frac{1}{2}$, wobei GH, HJ, JK, KG ihre an die Quadrate um $X, Y = 0, 0; 1, 0; 1, 1; 0, 1$ anstoßenden Seiten seien. Dann ist die Seite GK gleich der Seite des

Quadrates $\mathfrak{P}\left(\frac{\varrho}{2}, \frac{\sigma}{2}\right)$, die Seite GH kleiner als diese Seite. Der Grenzfall, daß $GH = GK$ wäre, würde nur eintreten, wenn A und A' beide in Ecken von $\mathfrak{P}(\varrho, \sigma)$ fielen, was dadurch ausgeschlossen ist, daß ξ, η nicht äquivalent der Form XY sein soll. Nunmehr erweist sich durch eine entsprechende Überlegung wie in 1°, daß der Inhalt des Quadrates $\mathfrak{P}\left(\frac{\varrho}{2}, \frac{\sigma}{2}\right)$ zusammen mit dem des Rechteckes $GHJK$ gleich 1 sein muß, woraus

$$(1) \quad \frac{1}{2} \varrho \sigma \leq 1 < 2 \cdot \frac{1}{2} \varrho \sigma$$

hervorgeht. Die Gleichung $\frac{1}{2} \varrho \sigma = 1$ hat statt, wenn die Lücken zum Fortfall kommen, also H mit G , J mit K zusammenfällt.

Es verhalte sich die Entfernung des Punktes A von der Geraden HJ zu der des Punktes L von dieser Geraden wie $1 : \kappa - 1$, so ist $1 \leq \kappa < 2$. Konstruieren wir nun das Quadrat $\mathfrak{P}\left(\frac{\kappa \varrho}{2}, \frac{\kappa \sigma}{2}\right)$, so wird es die Hälfte der

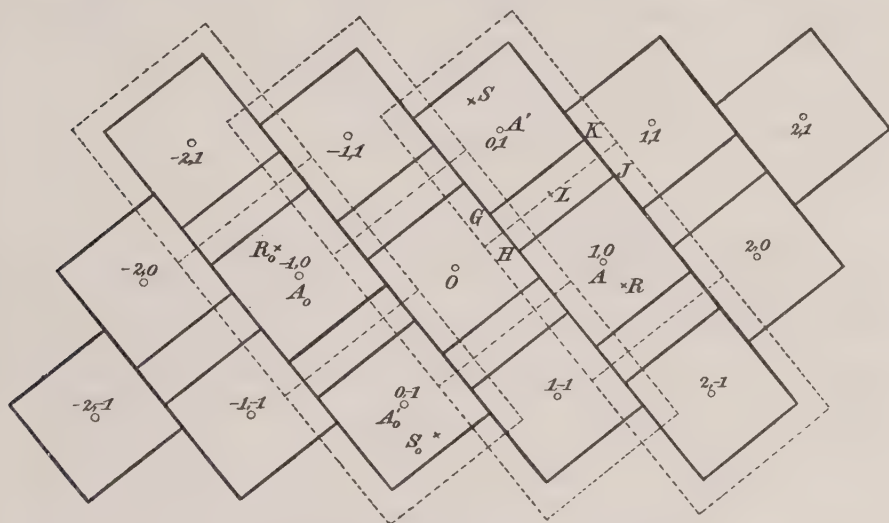


Fig. 6.

Lücke mit dem Mittelpunkte $X = -\frac{1}{2}$, $Y = \frac{1}{2}$ vollständig überdecken. Legen wir dann gleiche und parallel gestellte Quadrate um jeden Gitterpunkt als Mittelpunkt, so erfüllen daher diese Quadrate jedenfalls die ganze Ebene. Also wird derjenige Punkt, für den die Bestimmungsstücke ξ, η die gegebenen Werte $\xi = \xi_0$, $\eta = \eta_0$ haben, in wenigstens eines dieser Quadrate fallen müssen; alsdann gilt für den Gitterpunkt x, y , welcher der Mittelpunkt des betreffenden Quadrates ist,

$$(2) \quad \left| \frac{\xi - \xi_0}{\varrho} \right| + \left| \frac{\eta - \eta_0}{\sigma} \right| \leq \frac{\kappa}{2}.$$

Andererseits überdecken diese neuen Quadrate keinen Teil der Ebene mehr als *zweimal*, während noch gewisse Rechtecke mit den einzelnen Gitterpunkten als Mittelpunkten da sind, welche jedesmal nur je einem dieser Quadrate angehören. Infolgedessen muß der Inhalt des Quadrates $\mathfrak{P}\left(\frac{\kappa\rho}{2}, \frac{\kappa\sigma}{2}\right)$ kleiner als 2, also

$$(3) \quad \frac{1}{2} \kappa^2 \rho \sigma < 2$$

sein. Aus (2) und (3) ergibt sich wie oben

$$|(\xi - \xi_0)(\eta - \eta_0)| < \frac{1}{4}.$$

Damit ist der Satz V vollständig bewiesen. —

Wir bemerken noch folgendes: Aus $1 < \rho\sigma$ und $\kappa^2\rho\sigma < 4$ entnimmt man $\kappa^2 < 4$. Aus (2) erhält man daher $|\xi - \xi_0| < \rho$; nach § 2 ist aber stets $\rho < 2\lambda$. Hat nun die Kette zu ξ, η kein letztes Glied, ist also $\frac{-\beta}{\alpha}$ irrational, so konvergiert im Verlaufe ihrer Glieder die Größe λ nach Null. In solchem Falle kann man daher einen Gitterpunkt x, y bestimmen, wofür $|(\xi - \xi_0)(\eta - \eta_0)| < \frac{1}{4}$ ausfällt und zudem $|\xi - \xi_0|$ unter einer beliebig vorgeschriebenen positiven Größe liegt.

§ 5.

Es sei jetzt a eine beliebige reelle Größe, und wir setzen $\xi = x - ay$, $\eta = y$. Die Form $(x - ay)y$ ist nur dann äquivalent mit der Form XY , wenn a eine ganze Zahl ist, und nur dann äquivalent mit $\frac{1}{2}(X^2 - Y^2)$, wenn a gleich einer ganzen Zahl vermehrt um $\frac{1}{2}$ ist. Von diesen zwei Fällen wollen wir absehen. Die Kette zu ξ, η hat hier ein bestimmtes erstes Glied p_0, q_0 ; für dasselbe soll $\frac{p_0}{1} = \frac{q_0}{0}$ und > 0 sein, also hat man $p_0 = 1, q_0 = 0$. Für das zweite Kettenglied p_1, q_1 soll $p_0q_1 - q_0p_1 = \varepsilon_0 = 1$, also $q_1 = 1$ und $|(p_1 - aq_1)q_1| < \frac{1}{2}$ sein; also ist dann p_1 gleich der an der Größe a nächstgelegenen ganzen Zahl h_0 , für die $|h_0 - a| < \frac{1}{2}$ ist. Setzt man sodann in den Formeln (5) und (6) des § 3: $i = 1, t_k = 0$, so resultiert $\frac{p_k}{q_k} = h_0 - t_1$. Die in § 3 erhaltenen Resultate führen nunmehr zu folgendem Satze:

Satz VI. *Es sei a eine beliebige reelle Größe, jedoch weder eine ganze Zahl, noch gleich einer ganzen Zahl $+\frac{1}{2}$. Man bilde in folgender Weise eine Reihe von ganzzahligen Systemen p_i, q_i ($i = 0, 1, 2, \dots$).*

Zunächst sei $p_0 = 1$, $q_0 = 0$, sodann $q_1 = 1$ und $p_1 = h_0$ die an der Größe a nächstgelegene ganze Zahl so, daß $|h_0 - a| < \frac{1}{2}$ ist. Allgemein, wenn man bis zu einem Systeme p_i, q_i gelangt ist, wofür noch $p_i - a q_i \neq 0$ ausfällt, stelle man den Quotienten $\frac{p_{i-1} - a q_{i-1}}{p_i - a q_i}$ her. Es sei ϑ_i sein Vorzeichen und g_i die größte in dem absoluten Betrage dieses Quotienten enthaltene ganze Zahl, weiter $h_i = g_i$ oder $= g_i + 1$, je nachdem

$|((p_{i-1} - a q_{i-1}) - \vartheta_i g_i (p_i - a q_i))(g_i q_i - \vartheta_i q_{i-1})| < \text{oder} \geq \frac{1}{2}$
ist. Man setze sodann

$$p_{i+1} = h_i p_i - \vartheta_i p_{i-1}, \quad q_{i+1} = h_i q_i - \vartheta_i q_{i-1}.$$

Der auf diese Weise ermittelten Reihe von Systemen p_i, q_i kommen folgende Eigenschaften zu:

1°. Wenn a rational ist, bricht die Reihe mit einem gewissen System p_w, q_w ab, wofür $p_w - a q_w = 0$ ist. Wenn a irrational ist, so läßt sich die Reihe dieser Systeme unbegrenzt fortsetzen.

2°. Für je zwei aufeinanderfolgende Systeme gilt stets

$$p_i q_{i+1} - q_i p_{i+1} = \vartheta_1 \vartheta_2 \cdots \vartheta_i = \pm 1.$$

Die Zahlen p_i, q_i sind stets relativ prim.

3°. Man hat

$$\frac{p_k}{q_k} = h_0 - \frac{\vartheta_1}{|h_1|} - \frac{\vartheta_2}{|h_2|} - \cdots - \frac{\vartheta_{k-1}}{|h_{k-1}|}, \quad (k = 1, 2, \dots).$$

Dabei sind p_k und q_k selbst denjenigen Ausdrücken gleich, die bei der naturgemäßen Darstellung der rechten Seite als Quotient zweier ganzer Funktionen der h_i und ϑ_i den Zähler bzw. Nenner abgeben.

4°. Man hat

$$0 < q_1 < q_2 < q_3 \cdots, \\ \frac{1}{2} > |p_1 - a q_1| > |p_2 - a q_2| > |p_3 - a q_3|, \dots$$

Umsomehr nehmen die Beträge $\left| \frac{p_k}{q_k} - a \right|$ mit wachsendem Index ab. Wenn a irrational ist, konvergieren die Brüche $\frac{p_k}{q_k}$ mit wachsendem Index nach der Größe a .

5°. Für jedes der Systeme p_k, q_k ($k = 1, 2, \dots$) gilt

$$|(p_k - a q_k) q_k| < \frac{1}{2}.$$

Umgekehrt: Ist x, y irgendein System von relativ primen ganzen Zahlen, wofür $y > 0$ und

$$|(x - a y) y| < \frac{1}{2}$$

gilt, so findet sich das Zahlenpaar x, y stets unter den Systemen $p_k, q_k (k = 1, 2 \dots)$.

Aus dem Satze V entnimmt man noch: Sind b, c irgend zwei weitere reelle Größen, so kann man stets ganze Zahlen x, y finden, so daß

$$|(x - ay - b)(y - c)| < \frac{1}{4}$$

ist, und zwar, wenn a irrational ist, noch derart, daß zugleich $|x - ay - b|$ unter einer beliebig kleinen positiven Größe liegt.*)

Die hier definierte Kettenbruchentwicklung mit den Näherungsbrüchen $\frac{p_1}{q_1}, \frac{p_2}{q_2}, \dots$ zur Annäherung an die Größe a soll der *Diagonalkettenbruch* für a heißen, mit Rücksicht darauf, daß diese Näherungsbrüche zu den Parallelogrammen mit den Linien $x - ay = 0$ und $y = 0$ als *Diagonalen* in einer ganz entsprechenden Beziehung stehen, wie die Näherungsbrüche bei der gewöhnlichen Kettenbruchentwicklung

$$a = l_0 + \frac{1}{l_1} + \frac{1}{l_2} + \dots,$$

wo l_0 eine ganze Zahl, $l_1, l_2, \dots$ lauter positive Zahlen sind, zu den Parallelogrammen mit dem Nullpunkt als Mittelpunkt und mit Seiten parallel zu $x - ay = 0, y = 0$. Für diese *gewöhnliche* (auch sogenannte *regelmäßige* oder *normale*) Kettenbruchentwicklung würde ich alsdann den charakteristischeren Namen *Parallelkettenbruch* in Vorschlag bringen.

Nach einem bekannten Satze von Lagrange kommt jedes Zahlenpaar x, y , wobei x, y relativ prim sind, $y > 0$ und $\left| \frac{x}{y} - a \right| < \frac{1}{2y^2}$ ist, als Zähler und Nenner eines Näherungsbruches der gewöhnlichen Kettenbruchentwicklung für a vor. Danach erscheinen alle Näherungsbrüche des Diagonalkettenbruches für a auch unter den Näherungsbrüchen des Parallelkettenbruches für a und wird also, falls nicht beide Entwicklungen zusammenfallen, der Parallelkettenbruch stets langsamer konvergieren.

Der Diagonalkettenbruch für a möge mit $a = DK \left(\begin{smallmatrix} \vartheta_1, \vartheta_2, \dots \\ h_0, h_1, h_2, \dots \end{smallmatrix} \right)$, der Parallelkettenbruch für a mit $a = PK(l_0, l_1, l_2, \dots)$ bezeichnet

*) Tschebyscheff hat (in einem russisch geschriebenen Aufsätze in den Mémoires der Petersburger Akademie, T. X, Appendix 4, 1866; Oeuvres, T. I, p. 679) gezeigt, daß, wenn a, b reelle Größen sind, stets ganze Zahlen x, y existieren, wofür $|(x - ay - b)y| < 2$ ist. Hermite hat (Crelles Journal, Bd. 88, 1880; Oeuvres, T. III) bewiesen, daß der Ausdruck links durch ganze Zahlen x, y kleiner als $\sqrt{\frac{2}{27}}$ gemacht werden kann. Der Satz im Texte ergibt ein schärferes Resultat, weil $\frac{1}{4} < \sqrt{\frac{2}{27}}$ ist.

werden. Die schon vorhin angedeutete geometrische Auffassung der Parallelkettenbrüche, welche der hier gegebenen Ableitung der Diagonalkettenbrüche entspricht, führt leicht zu folgendem Satze:

Um aus der Reihe der Systeme $p_i, q_i (i=0, 1, 2, \dots)$ für den Diagonalkettenbruch von a die Reihe der Zähler und Nenner der sämtlichen Näherungsbrüche des Parallelkettenbruches von a zu erhalten, hat man nur die erstere Reihe in der Weise zu erweitern, daß, so oft eine Zahl $\vartheta_i = 1$ (für ein $i \geq 1$) ausfällt, zwischen die beiden Systeme p_{i-1}, q_{i-1} und p_i, q_i noch das neue System $p_i - p_{i-1}, q_i - q_{i-1}$ eingefügt wird.

Man leitet aus diesem Satze die nachstehende Regel ab, welche erlaubt, aus einem Diagonalkettenbruche $a = DK \left(\begin{smallmatrix} \vartheta_1, \vartheta_2, \dots \\ h_0, h_1, h_2, \dots \end{smallmatrix} \right)$ sogleich den Parallelkettenbruch $a = PK(l_0, l_1, l_2, \dots)$ zu entnehmen:

Aus der Reihe der Zahlen $h_0, h_1, h_2, \dots$ entsteht die Reihe der Zahlen $l_0, l_1, l_2, \dots$, indem an Stelle einer jeden Zahl h_i substituiert wird:

$$\begin{aligned} h_i, & \quad \text{wenn } \vartheta_i = -1, \quad \vartheta_{i+1} = -1, \\ h_i - 1, & \quad \text{wenn } \vartheta_i = 1, \quad \vartheta_{i+1} = -1, \\ h_i - 1, 1, & \quad \text{wenn } \vartheta_i = -1, \quad \vartheta_{i+1} = 1, \\ h_i - 2, 1, & \quad \text{wenn } \vartheta_i = 1, \quad \vartheta_{i+1} = 1 \end{aligned}$$

ist; dabei hat man sich noch $\vartheta_0 = -1$ zu denken und dann von dieser Vorschrift auch für $i=0$ Gebrauch zu machen.

Die Diagonalkettenbrüche sind hiernach nicht bloß wegen der einfacheren Charakterisierung ihrer Näherungsbrüche leichter zu handhaben, sie enthalten auch alle Einzelheiten über die darzustellenden Größen, welche die Parallelkettenbrüche erkennen lassen.

Beispiel: Der Parallelkettenbruch für $\sqrt{13}$ ist

$$PK(3, 1, 1; 1, 1, 6, 1, 1; 1, 1, 6, 1, 1; \dots)$$

mit den Näherungsbrüchen

$$\frac{3}{1}, \frac{4}{1}, \frac{7}{2}, \frac{11}{3}, \frac{18}{5}, \frac{119}{33}, \frac{137}{38}, \frac{256}{71}, \frac{393}{109}, \frac{649}{180}, \frac{4287}{1189}, \frac{4936}{1369}, \frac{9223}{2558}, \dots,$$

der Diagonalkettenbruch für $\sqrt{13}$ ist

$$DK \left(\begin{smallmatrix} +, -, +, +, -, +, +, \dots \\ 4, 2; 2, 8, 2; 2, 8, 2; \dots \end{smallmatrix} \right)$$

mit den Näherungsbrüchen

$$\frac{4}{1}, \frac{7}{2}, \frac{18}{5}, \frac{137}{38}, \frac{256}{71}, \frac{649}{180}, \frac{4936}{1369}, \frac{9223}{2558}, \dots$$

d. i. dem $2, 3; 5, 7, 8; \dots 5k, 5k+2, 5k+3; \dots$ ten Näherungsbruch der ersten Entwicklung.

Dagegen würde diejenige Kettenbruchentwicklung

$$K\left(\begin{array}{c} \vartheta_1, \vartheta_2, \dots \\ h_0, h_1, h_2, \dots \end{array}\right) = h_0 - \frac{\vartheta_1}{h_1} - \frac{\vartheta_2}{h_2} - \dots,$$

wobei $\vartheta_1 = \pm 1$, $\vartheta_2 = \pm 1, \dots$ und die $h_1, h_2, \dots$ positive ganze Zahlen sind derart, daß die Reste $-\frac{\vartheta_k}{h_k} - \frac{\vartheta_{k+1}}{h_{k+1}} - \dots$ stets zwischen $-\frac{1}{2}$ und $\frac{1}{2}$ liegen, für $\sqrt{13}$:

$$K\left(\begin{array}{ccccccc} 1, & 1, & -1, & 1, & 1, & -1, & 1, \dots \\ 4, & 3; & 2, & 7, & 3; & 2, & 7, & 3; \dots \end{array}\right)$$

sein mit den Näherungsbrüchen

$$\frac{4}{1}, \frac{11}{3}, \frac{18}{5}, \frac{137}{38}, \frac{393}{109}, \frac{649}{180}, \frac{4936}{1369}, \frac{14159}{3927}, \dots$$

d. i. dem $2, 4, 5, 7, 9; \dots 5k, 5k+2, 5k+4; \dots$ ten Näherungsbrüche des Parallelkettenbruches für $\sqrt{13}$. Also erfüllt bei dieser Art der Entwicklung einerseits nicht jeder Näherungsbruch $\frac{x}{y}$ die Bedingung

$$\left| \frac{x}{y} - \sqrt{13} \right| < \frac{1}{2y^2},$$

andererseits kommt nicht jeder Bruch $\frac{x}{y}$ mit dieser Eigenschaft unter den Näherungsbrüchen vor.

Der Diagonalkettenbruch für die Basis der natürlichen Logarithmen lautet:

$$e = DK\left(\begin{array}{cccccc} 1, & -1, & 1, & -1, & \dots, & 1, & -1, & \dots \\ 3, & 3, & 2, & 5, & 2, & \dots, & 2m+1, & 2, & \dots \end{array}\right).$$

Die Zähler und Nenner der Näherungsbrüche dieses Kettenbruches geben also die sämtlichen Auflösungen der Bedingung

$$-\frac{1}{2} < (x - ey)y < \frac{1}{2}$$

in relativ primen positiven ganzen Zahlen x, y .

§ 6.

1. Ein unendlicher Diagonalkettenbruch

$$(1) \quad DK\left(\begin{array}{c} \vartheta_1, \vartheta_2, \dots \\ h_0, h_1, h_2, \dots \end{array}\right)$$

für eine irrationale Größe a soll *periodisch* heißen, wenn es eine positive Zahl v gibt, so daß von einem gewissen Index j an stets

$$(2) \quad \vartheta_k = \vartheta_{k+v}, \quad h_k = h_{k+v} \quad (k = j, j+1, j+2, \dots)$$

ist. Das System der Werte

$$\begin{pmatrix} \vartheta_j, \vartheta_{j+1}, \dots, \vartheta_{j+v-1} \\ h_j, h_{j+1}, \dots, h_{j+v-1} \end{pmatrix}$$

heiße dann eine *Periode* des Kettenbruches.

Wir beweisen zunächst die folgende Tatsache:

Ist der Diagonalkettenbruch für eine irrationale GröÙe a periodisch, so ist a Wurzel einer quadratischen Gleichung mit rationalen Koeffizienten.

Wir verwenden für die Kette zu den Formen $\xi = x - ay$, $\eta = y$ die in § 3 eingeführten Bezeichnungen. Durch die Substitution

$$T_i = \begin{pmatrix} p_{i-1}, p_i \\ q_{i-1}, q_i \end{pmatrix}$$

geht ξ in $\varepsilon_{i-1}\lambda_{i-1}X_i + \varepsilon_i\lambda_iY_i$ über, man hat also die Formel der Komposition:

$$(1, -a) T_i = (\varepsilon_{i-1}\lambda_{i-1}, \varepsilon_i\lambda_i).$$

Aus § 3, Gleich. (2) leitet man allgemeiner die Regel

$$(\varepsilon_{i-1}\lambda_{i-1}, \varepsilon_i\lambda_i) \begin{pmatrix} 0, -\vartheta_i \\ 1, h_i \end{pmatrix} \begin{pmatrix} 0, -\vartheta_{i+1} \\ 1, h_{i+1} \end{pmatrix} \dots \begin{pmatrix} 0, -\vartheta_{k-1} \\ 1, h_{k-1} \end{pmatrix} = (\varepsilon_{k-1}\lambda_{k-1}, \varepsilon_k\lambda_k),$$

$$i < k$$

ab, und daraus entnimmt man insbesondere eine Beziehung:

$$\varepsilon_k\lambda_k = \varepsilon_{i-1}\lambda_{i-1}r_{i,k} + \varepsilon_i\lambda_i s_{i,k},$$

wo $r_{i,k}$, $s_{i,k}$ gewisse ganze Zahlen sind, die bloß von den Werten ϑ_i , h_i , ϑ_{i+1} , h_{i+1} , $\dots$, ϑ_{k-1} , h_{k-1} abhängen. Da mit unbegrenzt wachsendem k die GröÙe λ_k nach Null konvergiert, ersieht man danach, daß das Verhältnis $\frac{\varepsilon_{i-1}\lambda_{i-1}}{\varepsilon_i\lambda_i}$ durch die unendliche Reihe der GröÙen ϑ_k , h_k für die sämtlichen

Indizes $k \geq i$ vollständig bestimmt ist.

Wenn nun mit irgendeinem Werte v für alle Indizes $k \geq j$ die Beziehungen (2) statthaben, muß daher notwendig

$$\frac{\varepsilon_{j+v-1}\lambda_{j+v-1}}{\varepsilon_{j+v}\lambda_{j+v}} = \frac{\varepsilon_{j-1}\lambda_{j-1}}{\varepsilon_j\lambda_j}$$

sein. Setzen wir $\frac{\varepsilon_{j+v-1}\lambda_{j+v-1}}{\varepsilon_{j-1}\lambda_{j-1}} = \tau$, wobei $0 < |\tau| < 1$ sein wird, so folgt

$$\varepsilon_{j+v-1}\lambda_{j+v-1} = \tau\varepsilon_{j-1}\lambda_{j-1}, \quad \varepsilon_{j+v}\lambda_{j+v} = \tau\varepsilon_j\lambda_j.$$

Nun hat man

$$(1, -a) T_{j+v} = (\tau\varepsilon_{j-1}\lambda_{j-1}, \tau\varepsilon_j\lambda_j), \quad (\varepsilon_{j-1}\lambda_{j-1}, \varepsilon_j\lambda_j) T_j^{-1} = (1, -a);$$

daraus entsteht

$$(1, -a) T_{j+v} T_j^{-1} = (\tau, -\tau a).$$

Bezeichnet man mit $\begin{pmatrix} p, r \\ q, s \end{pmatrix}$ das Koeffizientenschema der Substitution $T_{j+\nu} T_j^{-1}$, so bedeutet diese letzte Formel

$$p - aq = \tau, \quad r - as = -\tau a.$$

Daraus erhält man

$$(r - as) + (p - aq) a = 0,$$

und hierin kann nicht $q = 0$ sein, denn sonst müßte $p = \tau$ sein, während τ keine ganze Zahl ist, da $|\tau|$ zwischen 0 und 1 fällt.

2. Wir beweisen jetzt die Umkehrung der eben festgestellten Tatsache:

Satz VII. *Ist eine irrationale GröÙe a Wurzel einer quadratischen Gleichung mit rationalen Koeffizienten, so ist der Diagonalkettenbruch für a stets periodisch.*

Beweis. Es sei

$$n_0 a^2 + n_1 a + n_2 = 0$$

diejenige Gleichung für a , in der n_0, n_1, n_2 relativ prime ganze Zahlen sind und $n_0 > 0$ ist. Man hat $a = \frac{-n_1 \pm \sqrt{n_1^2 - 4n_0 n_2}}{2n_0}$, so daß die ganze Zahl $D = n_1^2 - 4n_0 n_2$ positiv und wegen der vorausgesetzten Irrationalität von a jedenfalls nicht das Quadrat einer ganzen Zahl ist. Die zweite Wurzel jener Gleichung $\frac{-n_1 \mp \sqrt{D}}{2n_0}$ werde mit $\bar{a}$ bezeichnet.

Wir setzen

$$\xi = x - ay = \alpha x + \beta y, \quad \eta = \frac{1}{a - \bar{a}} (x - \bar{a}y) = \gamma x + \delta y, \quad \xi = y.$$

Dabei haben ξ, η ebenso wie ξ, ξ die Determinante 1, und gilt eine Beziehung:

$$\xi = \eta + b\xi, \quad b = -\frac{1}{a - \bar{a}} = \mp \frac{n_0}{\sqrt{D}}.$$

Sodann entsteht

$$\mp \sqrt{D} \xi \eta = f = n_0 x^2 + n_1 xy + n_2 y^2.$$

Wir betrachten nunmehr die Kette zu den Formen ξ, η . Diese Kette wird jedenfalls nach beiden Seiten unbegrenzt sein. Wir verwenden für diese Kette die in § 3 eingeführten Bezeichnungen. Außerdem mögen ξ_i, η_i die Ausdrücke bedeuten, in welche ξ, η durch die Substitution

$$(T_i) \quad x = p_{i-1} X_i + p_i Y_i, \quad y = q_{i-1} X_i + q_i Y_i$$

übergehen, und es bedeute T_i das quadratische Schema $\begin{pmatrix} \varepsilon_{i-1} \lambda_{i-1} & \varepsilon_i \lambda_i \\ \mu_{i-1} & \mu_i \end{pmatrix}$ der Koeffizienten von ξ_i, η_i , wobei dann $\begin{pmatrix} \alpha & \beta \\ \gamma & \delta \end{pmatrix} T_i = T_i$ gilt.

Durch die ganzzahlige Substitution T_i , deren Determinante ± 1 ist, geht die Form $f = \pm \sqrt{D} \xi \eta$ in eine Form

$$f_i = \pm \sqrt{D} \xi_i \eta_i = N_0 X_i^2 + N_1 X_i Y_i + N_2 Y_i^2$$

über. Dabei sind N_0, N_1, N_2 ganze rationale Zahlen und folgt $N_1^2 - 4N_0N_2 = D$; da D nicht eine Quadratzahl ist, hat man gewiß $N_0 \neq 0, N_2 \neq 0$. Zudem bestehen dabei nach § 2, 3 diese Beziehungen:

entweder

$$\frac{N_2}{N_0} < 0, \quad \frac{N_1}{N_0} > 0, \quad |N_0| < \frac{1}{2}\sqrt{D}, \quad |N_2| < \frac{1}{2}\sqrt{D}, \quad |N_1| < \sqrt{D}$$

oder

$$\frac{N_2}{N_0} > 0, \quad \frac{N_1}{N_0} > 0, \quad |N_0| < \frac{1}{2}\sqrt{D}, \quad |N_2| < \frac{1}{2}\sqrt{D}, \quad \sqrt{D} < |N_1| < \frac{3}{2}\sqrt{D},$$

$$|N_1 - N_0 - N_2| > \frac{1}{2}\sqrt{D}.$$

In jedem Falle kommen hiernach für die ganzzahligen Koeffizienten N_0, N_1, N_2 einer Form f_i von vornherein nur eine endliche Anzahl von möglichen Wertsystemen in Betracht. Man wird also jedenfalls irgend zwei Indizes $i = j$ und $i = j + v$, wo $v > 0$ ist, finden können, für welche die beiden Formen f_j und f_{j+v} in den Koeffizienten N_0, N_1, N_2 übereinstimmen.

Ersetzt man X_{j+v}, Y_{j+v} in den Formen ξ_{j+v}, η_{j+v} durch die Zeichen X_j, Y_j der Variablen in ξ_j, η_j , so werden dadurch Beziehungen

$$\xi_{j+v} = A\xi_j + B\eta_j, \quad \eta_{j+v} = \Gamma\xi_j + \Delta\eta_j$$

hergestellt, wobei die Koeffizienten A, B, Γ, Δ durch $T_{j+v} = \begin{pmatrix} A, & B \\ \Gamma, & \Delta \end{pmatrix} T_j$ be-

bestimmt sind, und vermöge dieser Beziehungen muß dann $\xi_{j+v}\eta_{j+v} = \xi_j\eta_j$ entstehen. Vergleicht man die Koeffizienten von $\xi_j^2, \eta_j^2, \xi_j\eta_j$ auf beiden Seiten dieser Gleichung, so folgt $A\Gamma = 0, B\Delta = 0, A\Delta + B\Gamma = 1$. Danach muß entweder

$$(3) \quad B = 0, \quad \Gamma = 0, \quad A = \frac{1}{\Delta} = \tau; \quad \xi_{j+v} = \tau\xi_j, \quad \eta_{j+v} = \frac{1}{\tau}\eta_j$$

oder

$$(3^*) \quad A = 0, \quad \Delta = 0, \quad B = \frac{1}{\Gamma} = \tau; \quad \xi_{j+v} = \tau\eta_j, \quad \eta_{j+v} = \frac{1}{\tau}\xi_j$$

mit einem von Null verschiedenen Faktor τ sein. Die zweite Art von Beziehungen aber ist unmöglich, weil in der Form η_j der zweite Koeffizient einen größeren absoluten Betrag hat als der erste Koeffizient, in der Form ξ_{j+v} aber das Entgegengesetzte statthaben muß. Also gelten notwendig die Beziehungen (3) und die Vergleichung der Koeffizienten in η_j und η_{j+v} zeigt noch, daß τ positiv und < 1 ist.

Die Gleichungen (3) ergeben

$$T_{j+v} = \begin{pmatrix} \tau, & 0 \\ 0, & \frac{1}{\tau} \end{pmatrix} T_j, \quad \begin{pmatrix} \alpha, & \beta \\ \gamma, & \delta \end{pmatrix} T_{j+v} T_j^{-1} = \begin{pmatrix} \tau, & 0 \\ 0, & \frac{1}{\tau} \end{pmatrix} \begin{pmatrix} \alpha, & \beta \\ \gamma, & \delta \end{pmatrix}.$$

Ist $\begin{pmatrix} p, & r \\ q, & s \end{pmatrix}$ das Koeffizientenschema der Substitution $T_{j+v} T_j^{-1}$, wobei p, q, r, s ganze Zahlen sind und $ps - qr = \pm 1$ ist, so verwandeln sich hiernach ξ, η , wenn man darin x, y durch $px + ry, qx + sy$ ersetzt, in die Ausdrücke $\tau \xi, \frac{1}{\tau} \eta$. Diese zwei Ausdrücke ergeben dasselbe Produkt wie ξ und η . Durch die umgekehrte Substitution werden ξ, η in $\frac{1}{\tau} \xi, \tau \eta$ übergehen. Hat man nun einen Gitterpunkt x, y , für den x, y relativ prim sind und $\xi = \varepsilon \lambda, \eta = \mu$ ($\mu > 0, \lambda > 0, \varepsilon = \pm 1$), $\lambda \mu < \frac{1}{2}$ ist, so existiert dann also ein anderer Gitterpunkt, für den ebenfalls x, y relativ prim sind und $\xi = \tau \varepsilon \lambda, \eta = \frac{1}{\tau} \mu$ ist, und ferner ein Gitterpunkt, für den x, y relativ prim sind und $\xi = \frac{1}{\tau} \varepsilon \lambda, \eta = \tau \mu$ ist; und diese zwei weiteren Gitterpunkte müssen dann ebenso wie der erste als Glieder der Kette zu ξ, η auftreten.

Nach (3) haben wir $\varepsilon_{j+v} \lambda_{j+v} = \tau \varepsilon_j \lambda_j, \mu_{j+v} = \frac{1}{\tau} \mu_j$. Nun müssen weiter die Punkte $\xi = \varepsilon_{j+v+1} \lambda_{j+v+1}, \eta = \mu_{j+v+1}$ und $\xi = \tau \varepsilon_{j+1} \lambda_{j+1}, \eta = \frac{1}{\tau} \mu_{j+1}$, welche beide als Glieder der Kette auftreten, identisch sein. Denn hätte man $\mu_{j+v+1} > \frac{1}{\tau} \mu_{j+1}$, so würde der Punkt $\xi = \tau \varepsilon_{j+1} \lambda_{j+1}, \eta = \frac{1}{\tau} \mu_{j+1}$ ein Kettenglied sein, für das $\mu_{j+v} < \eta < \mu_{j+v+1}$ wäre, während $\eta = \mu_{j+v}$ und $\eta = \mu_{j+v+1}$ ja zwei aufeinanderfolgenden Kettengliedern entsprechen. Hätte man dagegen $\mu_{j+v+1} < \frac{1}{\tau} \mu_{j+1}$, so würde $\xi = \frac{1}{\tau} \varepsilon_{j+v+1} \lambda_{j+v+1}, \eta = \tau \mu_{j+v+1}$ ein Kettenglied sein, wofür $\mu_j < \eta < \mu_{j+1}$ wäre, was ebenfalls nicht möglich ist. Also muß $\mu_{j+v+1} = \frac{1}{\tau} \mu_{j+1}$ sein und müssen sodann jene zwei Punkte zusammenfallen.

Auf dieselbe Art erschließt man sukzessive weiter

$$(4) \quad \varepsilon_{k+v} \lambda_{k+v} = \tau \varepsilon_k \lambda_k, \quad \mu_{k+v} = \frac{1}{\tau} \mu_k$$

für $k = j + 2, j + 3, \dots$. Sodann kann man in der Reihe der Indizes rückwärts gehen und diese Beziehungen (4), welche auch für $k = j - 1$ bereits durch (3) feststehen, nacheinander für $k = j - 2, j - 3, \dots$ erhalten. Also gelten diese Beziehungen (4) überhaupt für jeden möglichen

Index k . Man hat nunmehr für jeden Index i : $\tau_{i+v} = \begin{pmatrix} \tau, & 0 \\ 0, & \frac{1}{\tau} \end{pmatrix} \tau_i$; andererseits gilt nach § 3, (2) die Regel $\tau_{i+1} = \tau_i \begin{pmatrix} 0, & -\vartheta_i \\ 1, & h_i \end{pmatrix}$. Wendet man

die erstere Formel für zwei aufeinanderfolgende Indizes $i = k$, $i = k + 1$ an und hernach die letztere für $i = k$ und $i = k + v$, so folgt zuerst $T_{k+v}^{-1} T_{k+v+1} = T_k^{-1} T_{k+1}$ und sodann:

$$(5) \quad \vartheta_{k+v} = \vartheta_k, \quad h_{k+v} = h_k \quad (k = \dots - 2, -1, 0, 1, 2, \dots).$$

Nach diesen Beziehungen (5) kann die Kette zu ξ, η als *vollkommen periodisch* bezeichnet werden.

Wir ziehen endlich auch die Form $\xi = y$ heran. Für ein jedes Glied $x = p_i$, $y = q_i$ der Kette zu den Formen ξ, η gilt $\eta > 0$ und $|\xi\eta| < \frac{1}{2}$. Gleichzeitig ist dabei $\sqrt{D}|\xi\eta|$ stets eine rationale ganze Zahl. Indem diese Zahl $< \frac{1}{2}\sqrt{D}$ ist, kann sie nicht die größte in $\frac{1}{2}\sqrt{D}$ enthaltene ganze Zahl überschreiten. Wir setzen letztere ganze Zahl $\left[\frac{1}{2}\sqrt{D}\right] = \frac{1}{2}\sqrt{D} - d$, dabei wird $0 < d < 1$. Aus $\sqrt{D}|\xi\eta| \leq \frac{1}{2}\sqrt{D} - d$ folgt mit Rücksicht auf $\xi = y = \eta + b\xi$:

$$|\xi\xi| \leq \frac{1}{2} - \frac{d}{\sqrt{D}} + |b\xi^2|, \quad \xi \geq \eta - |b\xi|.$$

Nun nimmt, wenn man die Reihe der Kettenglieder zu ξ, η durchläuft, darin sowohl $|\xi|$ wie $\left|\frac{\xi}{\eta}\right|$ beständig ab. Geht man soweit in dieser Reihe, daß $|b\xi^2| < \frac{d}{\sqrt{D}}$ und $\left|\frac{\xi}{\eta}\right| < \frac{1}{|b|}$ ist, so hat man dann also $|\xi\xi| < \frac{1}{2}$, $\xi > 0$ und müssen daher die betreffenden Systeme $x = p_i$, $y = q_i$, die man nun antrifft, sämtlich auch Glieder der Kette zu den Formen ξ, ξ sein.

Ist umgekehrt x, y ein Glied der Kette zu den Formen ξ, ξ , so ist dafür $\xi > 0$ und $|\xi\xi| < \frac{1}{2}$; daraus folgt

$$\sqrt{D}|\xi\eta| < \frac{1}{2}\sqrt{D} + |\sqrt{D}b\xi^2|, \quad \eta \geq \xi - |b\xi|.$$

Geht man nun in der Reihe der Kettenglieder zu ξ, ξ so weit, daß $|\sqrt{D}b\xi^2| \leq 1 - d$ und $\left|\frac{\xi}{\xi}\right| < \frac{1}{|b|}$ ist, so folgt hier $\eta > 0$ und andererseits $\sqrt{D}|\xi\eta| < \left[\frac{1}{2}\sqrt{D}\right] + 1$. Da aber $\sqrt{D}|\xi\eta|$ hierbei stets eine ganze rationale Zahl wird, muß dann überhaupt $\sqrt{D}|\xi\eta| \leq \left[\frac{1}{2}\sqrt{D}\right] < \frac{1}{2}\sqrt{D}$, mithin $|\xi\eta| < \frac{1}{2}$ sein. Danach findet sich alsdann das betreffende System x, y stets auch unter den Gliedern der Kette zu ξ, η .

Auf diese Weise stimmen überhaupt die zu ξ, η und die zu ξ, ξ gehörige Kette von gewissen zwei Gliedern an in dem ganzen weiteren Verlaufe ihrer Kettenglieder völlig miteinander überein. Da nun die einzelnen Werte ϑ_i, h_i in einer Kette jedesmal aus drei aufeinanderfolgenden

Kettengliedern abzuleiten sind, so werden sich danach auch die Relationen (5) von einem gewissen Index an auf die Kette zu $\xi = x - ay$, $\xi = y$ übertragen, d. h. eben der Diagonalkettenbruch für die Größe a ist periodisch.

In entsprechender Beziehung wie der Diagonalkettenbruch für a zu der Kette steht, die zu den Formen ξ, η gehört, steht der Diagonalkettenbruch für die konjugierte algebraische Zahl $\bar{a}$ zu der Kette, die zu den Formen $(a - \bar{a})\eta, -\frac{1}{a - \bar{a}}\xi$ oder, was auf dasselbe hinauskommt, zu $\eta, -\xi$ gehört. Der einfache Zusammenhang der Kette zu ξ, η und der Kette zu $\eta, -\xi$ ist am Schlusse von § 3 erörtert; aus der dort angegebenen Beziehung (7) erhellt nun, daß, wenn $\begin{pmatrix} \vartheta_j, \vartheta_{j+1}, \dots, \vartheta_{j+v-1} \\ h_j, h_{j+1}, \dots, h_{j+v-1} \end{pmatrix}$ eine Periode des Diagonalkettenbruches für a ist, $\begin{pmatrix} \vartheta_j, \vartheta_{j+v-1}, \dots, \vartheta_{j+1} \\ h_{j+v-1}, h_{j+v-2}, \dots, h_j \end{pmatrix}$ eine Periode des Diagonalkettenbruches für $\bar{a}$ sein wird.

Zürich, den 31. Dezember 1899.

XVII.

Quelques nouveaux théorèmes sur l'approximation des quantités a l'aide de nombres rationnels.

(Bulletin des Sciences mathématiques, 2^e série, t. XXV, pp. 72—76.)

Dans plusieurs de ses Mémoires classiques sur la théorie des nombres M. Hermite s'est occupé de la recherche des minima de formes algébriques pour des valeurs entières des variables. Cette recherche l'a conduit à certaines approximations, dont il a mis parfaitement en évidence le *caractère algébrique*. Mais pour la plupart des inégalités que l'on obtient dans ce domaine, il y a encore une grande difficulté à déterminer les *limites les plus étroites* des inégalités et les *formes extrêmes* auxquelles en chaque cas ces limites sont rattachées. Dans quelques cas particulièrement simples je donnerai ici de telles limites précises; pour les démonstrations, un peu longues, je dois renvoyer le lecteur au deuxième cahier de ma *Géométrie des nombres*, qui paraîtra prochainement chez B. G. Teubner.*)

I.

1. *Théorème.* — Soient $\varphi, \chi, \psi, \omega$ quatre formes linéaires à trois variables x, y, z , à coefficients réels quelconques et de sorte que l'on ait

$$\varphi + \chi + \psi + \omega = 0.$$

Supposons que le déterminant de trois de ces formes soit toujours différent de zéro et désignons sa valeur absolue par $4D$.

Alors il existe toujours trois nombres entiers x, y, z , qui ne sont pas tous égaux à zéro et de sorte que toutes les quatre formes $\varphi, \chi, \psi, \omega$ soient en valeur absolue moindres que

$$\sqrt[3]{\frac{4D}{1 - \left(\frac{2}{3}\right)^3}}.$$

La limite $d = \sqrt[3]{\frac{108}{19}D}$ donnée ici est précise. En général, on peut aussi trouver des nombres entiers x, y, z différents du système 0, 0, 0 et tels que les valeurs absolues de $\varphi, \chi, \psi, \omega$ soient toutes $< d$. Mais

*) Vgl. Diophantische Approximationen, Kap. III und insbesondere Kap. VI. (Anm. d. Herausg.)

il y a un cas d'exception, où il sera impossible de satisfaire à cette autre condition, c'est lorsqu'il existe une substitution linéaire à coefficients entiers et à déterminant ± 1 , transformant les formes $\varphi, \chi, \psi, \omega$, abstraction faite de l'ordre, en

$$d\left(X - \frac{2}{3}Y\right), \quad d\left(Y - \frac{2}{3}Z\right), \quad d\left(-\frac{2}{3}X + Z\right), \quad -\frac{1}{3}d(X + Y + Z).$$

2. Le théorème que nous venons d'énoncer est susceptible d'une interprétation géométrique, qui peut-être trouvera une application dans la cristallographie:

Imaginons un système d'octaèdres, tous congruents entre eux et parallèlement orientés, en nombre infini et placés dans l'espace de manière que deux de ces corps n'aient jamais une partie commune et que leurs centres de gravité forment un réseau parallélépipédique de points, comme Bravais l'a considéré dans sa théorie de la structure des cristaux. (Tous les octaèdres peuvent alors être déduits de l'un quelconque d'entre eux par le même système de translations). Il y aura une situation où les lacunes laissées par les octaèdres seront les plus étroites, ou, ce qui est la même chose, l'espace occupé par les corps sera le plus grand possible; c'est la situation où chacun des octaèdres a au moins un point commun avec quatorze autres de ces corps, et alors le rapport de l'espace occupé par les octaèdres à l'espace laissé libre par ces corps sera de 18 à 1.

3. Soient ξ, η, ζ trois formes en x, y, z à coefficients réels et à déterminant $\pm D$, où $D > 0$. On pourra prendre dans le théorème du n° 1: $\varphi = -\xi + \eta + \zeta, \quad \chi = \xi - \eta + \zeta, \quad \psi = \xi + \eta - \zeta, \quad \omega = -\xi - \eta - \zeta$. Il existe alors des nombres entiers x, y, z différents du système 0, 0, 0 et de sorte que l'on ait

$$|\xi| + |\eta| + |\zeta| \leq \sqrt[3]{\frac{108}{19}} D.$$

Dans le cas limite, pour lequel le signe $=$ est réservé, on peut toujours faire en sorte que l'on n'ait pas à la fois $|\xi| = |\eta| = |\zeta|$.

On en tire alors

$$|\xi\eta\zeta| < \frac{4}{19} D.$$

Or on doit remarquer que, dans cette dernière inégalité, le facteur $\frac{4}{19}$ pourrait encore être remplacé par un nombre plus petit.

4. D'autre part, on peut aussi prendre dans le n° 1:

$$\frac{\varphi}{\sqrt[3]{2}}, \quad \frac{\chi}{\sqrt[3]{2}}, \quad \frac{\psi}{\sqrt[3]{2}}, \quad \frac{\omega}{\sqrt[3]{2}} = \xi + \zeta, \quad -\xi + \zeta, \quad \eta - \zeta, \quad -\eta - \zeta.$$

On voit alors qu'il existe des nombres entiers x, y, z différents du système 0, 0, 0 de sorte que l'on ait

$$|\xi| + |\zeta| \leq \sqrt[3]{\frac{54}{19}} D, \quad |\eta| + |\zeta| \leq \sqrt[3]{\frac{54}{19}} D.$$

Dans le cas limite où l'on doit prendre ici un des deux signes =, on peut toujours faire en sorte que l'on n'ait ni $\pm \xi = \zeta$ ni $\pm \eta = \zeta$.

En faisant usage alors des inégalités

$$\left| \left(\frac{\xi}{2} \right)^2 \zeta \right| \leq \left(\frac{2 \left| \frac{\xi}{2} \right| + |\zeta|}{3} \right)^3, \quad \left| \left(\frac{\eta}{2} \right)^2 \zeta \right| \leq \left(\frac{2 \left| \frac{\eta}{2} \right| + |\zeta|}{3} \right)^3,$$

on trouve

$$|\xi^2 \zeta| < \frac{8}{19} D, \quad |\eta^2 \zeta| < \frac{8}{19} D.$$

Dans ces dernières inégalités, la limite n'est plus précise.

5. Soient a, b deux quantités réelles quelconques. Posons, dans les inégalités du n° 4:

$$\xi = x - az, \quad \eta = y - bz, \quad \zeta = \frac{z}{t^3},$$

t étant un paramètre positif. On voit qu'on peut trouver des nombres entiers x, y, z , parmi lesquels z est positif, et de sorte que $x - az, y - bz$ soient en valeur absolue plus petites qu'une quantité ε donnée à volonté, et que dans cette approximation on ait en même temps

$$\left| \frac{x}{z} - a \right| < \sqrt[3]{\frac{8}{19}} \frac{1}{z^{\frac{2}{3}}}, \quad \left| \frac{y}{z} - b \right| < \sqrt[3]{\frac{8}{19}} \frac{1}{z^{\frac{2}{3}}}.$$

La constante $\sqrt[3]{\frac{8}{19}} = 0,648 \dots$, qui entre ici, est plus petite que $\frac{2}{3}$.

II.

6. *Théorème.* — Soient $\xi = \alpha x + \beta y, \eta = \gamma x + \delta y$ deux formes linéaires à coefficients complexes quelconques et soit $D = |\alpha\delta - \beta\gamma| > 0$. On peut toujours trouver, dans le corps algébrique de $i = \sqrt{-1}$, des nombres entiers complexes x, y différents du système 0, 0 de sorte que l'on ait

$$|\xi| \leq \sqrt{\frac{\sqrt{3}+1}{\sqrt{6}}} D, \quad |\eta| \leq \sqrt{\frac{\sqrt{3}+1}{\sqrt{6}}} D.$$

La limite $d = \sqrt{\frac{\sqrt{3}+1}{\sqrt{6}}} D$ donnée ici est précise. En général, on aura aussi des nombres entiers complexes $x, y \neq 0, 0$ de sorte que $|\xi| < d, |\eta| < d$. Mais dans un seul cas il sera impossible de satisfaire à ces autres inégalités; c'est lorsqu'il existe une substitution

$$x = pX + rY, \quad y = qX + sY,$$

où p, q, r, s sont des nombres entiers dans le corps de i et le détermi-

nant $ps - qr = \pm 1$ ou $\pm i$, de sorte que, par cette substitution, ξ, η soient transformées en

$$\lambda d \left\{ X + \left[\frac{1}{2} - i \left(1 - \frac{\sqrt{3}}{2} \right) \right] Y \right\}, \quad \mu d \left\{ \left[\frac{i}{2} + \left(1 - \frac{\sqrt{3}}{2} \right) \right] X + Y \right\},$$

λ et μ étant des quantités dont la valeur absolue est égale à 1.

7. *Théorème.* — Soient $\xi = \alpha x + \beta y$, $\eta = \gamma x + \delta y$ deux formes linéaires à coefficients complexes quelconques et soit $D = |\alpha\delta - \beta\gamma| > 0$. On peut toujours trouver, dans le corps algébrique de $j = \frac{-1 + \sqrt{-3}}{2}$, des nombres entiers complexes x, y différents du système 0, 0 de sorte que l'on ait

$$|\xi| \leq \sqrt{D}, \quad |\eta| \leq \sqrt{D}.$$

Cette limite $d = \sqrt{D}$ est ici précise. En général, il y aura aussi des nombres entiers $x, y \neq 0, 0$ de sorte que $|\xi|$ et $|\eta|$ soient $< d$, excepté dans le cas où, dans le corps de j , il existe une substitution linéaire à coefficients entiers et à déterminant égal à une unité du corps, à l'aide de laquelle ξ, η , abstraction faite de l'ordre, se changent en $\lambda d X, \mu d (\tau X + Y)$, λ et μ étant des quantités dont la valeur absolue est égale à 1, et τ une quantité complexe quelconque.

On remarquera ce fait intéressant que, de ces deux théorèmes correspondants, l'un, qui est relatif au corps algébrique de la *troisième racine de l'unité*, est beaucoup plus simple que l'autre, relatif au corps algébrique de la *quatrième racine de l'unité*.

8. Soit a une quantité complexe quelconque. En posant

$$\xi = x - ay, \quad \eta = \frac{y}{t},$$

t étant un paramètre réel quelconque > 1 , on voit que, dans le corps de $\frac{-1 + \sqrt{-3}}{2}$ (mais pas dans le corps de $\sqrt{-1}$), il y aura toujours des nombres entiers complexes x, y tels que

$$0 < |y| \leq t, \quad |x - ay| < \frac{1}{t},$$

d'où l'on tire encore

$$|(x - ay)y| < 1.$$

XVIII.

Über periodische Approximationen algebraischer Zahlen.

(Acta Mathematica, Band 26, (Niels Henrik Abel in memoriam), S. 333—351.)

Abel sagt an einer Stelle (Oeuvres, T. II, p. 217) mit Bezug auf das Problem der algebraischen Auflösung der Gleichungen: „Au lieu de demander une relation dont on ne sait pas si elle existe ou non, il faut demander si une telle relation est en effet possible“. Eben diese Weisung befolgend, können wir auch einer anderen, noch unerledigten Aufgabe auf dem mannigfaltigen Gebiete der Auflösung der Gleichungen näherzutreten versuchen. Wir wollen hier die Frage behandeln:

Welche algebraische Zahlen besitzen analoge periodische Approximationen, wie sie die reellen algebraischen Zahlen zweiten Grades vermöge der Periodizität ihrer Entwicklungen in gewöhnliche Kettenbrüche aufweisen.

§ 1. Periodische Substitutionenketten.

1. Es sei α eine beliebige Größe und es werde $l = 1$ oder $= 2$ gesetzt, je nachdem α reell oder komplex ist. Wenn α eine *algebraische Zahl* n^{ten} Grades, d. h. eine Wurzel einer im Bereiche der rationalen Zahlen irreduziblen Gleichung n^{ten} Grades ist, so kann der Ausdruck

$$\xi = x_1 + \alpha x_2 + \cdots + \alpha^{n-1} x_n$$

für ganze rationale Zahlen $x_1, x_2, \dots, x_n$, die nicht sämtlich Null sind, niemals verschwinden, aber, wofern $n > l$ ist, wohl dem Werte Null beliebig nahe kommen. Wir machen hier stets die Annahme $n > l$. Über die Annäherungen dieser Form ξ an Null gelten dann, wie ich in dem Aufsätze „*Ein Kriterium für die algebraischen Zahlen*“ (Göttinger Nachrichten v. 11. Febr. 1899; diese Ges. Abhandlungen, Bd. I, S. 293—315) gezeigt habe, die folgenden Sätze:

Wir können zur Zahl α in bezug auf jede beliebige reelle Größe $r \geq 1$ stets eine Substitution

$$S) \quad x_h = s_h^{(1)} y_1 + s_h^{(2)} y_2 + \cdots + s_h^{(n)} y_n \quad (h = 1, 2, \dots, n)$$

mit folgenden Eigenschaften konstruieren:

a) Alle Koeffizienten $s_h^{(k)}$ sind ganze rationale Zahlen, und gehen die Quotienten $\frac{s_h^{(k)}}{r}$ dem Betrage nach nicht über eine gewisse, von r unabhängige GröÙe hinaus.

b) Die Determinante von S ist $\neq 0$ und liegt dem Betrage nach unter einer gewissen, von r nicht abhängigen Grenze.

c) Geht ξ durch S in

$$\varphi = \varrho_1 y_1 + \varrho_2 y_2 + \cdots + \varrho_n y_n$$

über, so liegen die Beträge von

$$\varrho_1 r^{\frac{n-l}{l}}, \varrho_2 r^{\frac{n-l}{l}}, \dots, \varrho_n r^{\frac{n-l}{l}}$$

sämtlich unter einer gewissen, von r nicht abhängigen Grenze.

d) Für die Verhältnisse $\varrho_1 : \varrho_2 : \dots : \varrho_n$ kommen von vornherein nur eine endliche Anzahl verschiedener Systeme in Betracht, die von r nicht abhängen.

Diesen Bedingungen wird z. B., wie in jener Arbeit ausgeführt ist, stets genügt, wenn wir S unter allen denjenigen Substitutionen, für welche die Koeffizienten $s_h^{(k)}$ lauter Zahlen aus der Reihe $0, \pm 1, \pm 2, \dots, \pm [r]$ sind und die Determinante $\neq 0$ ist, derart auswählen, daß dabei zunächst $|\varrho_1|$, nächst dem $|\varrho_2|$, $\dots$ endlich $|\varrho_n|$ möglichst klein werden. Dabei fällt dann die Determinante von S dem Betrage nach sicher stets $\leq n!$ aus.

Durch die Substitution S erlangen wir zugleich gewisse rationale Approximationen für alle Zahlen des Körpers von α , wenn α reell ist, bzw., wenn α komplex ist, für alle reellen Zahlen des Körpers, der aus dem Körper von α und dem dazu konjugiert imaginären Körper zusammengesetzt ist.

Umgekehrt gilt der Satz: Die GröÙe α ist, ($n > l$ angenommen), notwendig eine algebraische Zahl n^{ten} Grades, wenn für sie in bezug auf jede reelle GröÙe $r \geq 1$ stets eine den Bedingungen a), b), c), d) entsprechende Substitution S hergestellt werden kann.

2. Wir denken uns weiterhin α stets als eine algebraische Zahl n^{ten} Grades und $n > l$. Nehmen wir nun eine unbegrenzte Reihe wachsender Zahlen $r_1 \geq 1, r_2, r_3, \dots$ an und konstruieren wir in der eben erörterten Weise zu diesen Zahlen Substitutionen $S_1, S_2, S_3, \dots$. Eine derartige Substitutionenkette für die Zahl α soll periodisch heißen, wenn die daraus vermittelte Kompositionsformeln

$$S_2 = S_1 Q_1, \quad S_3 = S_2 Q_2, \dots, S_{j+1} = S_j Q_j, \dots$$

hergeleitete Reihe von Substitutionen $Q_1, Q_2, Q_3, \dots$, abgesehen von einer endlichen Anzahl von Gliedern am Anfange, in periodischer Wiederholung ein und derselben endlichen Folge von Substitutionen besteht, wenn also ein

Index j_0 und eine positive Zahl p_0 angebar sind, so daß für jeden beliebigen Index $j \geq j_0$ stets $Q_j = Q_{j+p_0}$ ist.

Wir fragen nach dem Charakter derjenigen algebraischen Zahlen α , für welche periodische Substitutionsketten existieren.

3. Ist die Kette $S_1, S_2, S_3, \dots$ für α periodisch, so erhalten wir mit den soeben eingeführten Bezeichnungen

$$Q_j = S_j^{-1} S_{j+1}, \quad S_j^{-1} S_{j+1} = S_{j+p_0}^{-1} S_{j+p_0+1}, \quad j \geq j_0,$$

also

$$S_{j+p_0} S_j^{-1} = S_{j+p_0+1} S_{j+1}^{-1},$$

wenn $j \geq j_0$ ist. Setzen wir $S_{j_0+p_0} S_{j_0}^{-1} = P_0$, so folgt daraus allgemein

$$S_{j+p_0} = P_0 S_j, \quad S_{j+f p_0} = P_0^f S_j$$

für jeden Index $j \geq j_0$ und jeden Exponenten $f = 1, 2, 3, \dots$

Es bedeute φ_j die Linearform, in welche ξ durch S_j übergeht. Unter den unendlich vielen Substitutionen $S_{j_0+f p_0}$ für $f = 0, 1, 2, \dots$ werden wir, da nach b) für ihre Determinanten und weiter nach d) für die Verhältnisse der Koeffizienten $\varphi_1 : \varphi_2 : \dots : \varphi_n$ in den zugehörigen Formen $\varphi_{j_0+f p_0}$ nur eine endliche Anzahl von Wertsystemen in Betracht kommen, jedenfalls irgend zwei Substitutionen $S_{j_0+c p_0} = S$ und $S_{j_0+d p_0} = T$ ($d > c$) in solcher Weise auffinden können, daß erstens $TS^{-1} = P = P_0^{d-c}$ eine ganzzahlige Substitution mit der Determinante 1 wird und zudem zweitens in den beiden Formen $\varphi_{j_0+c p_0} = \varphi$ und $\varphi_{j_0+d p_0} = \psi$ die n Koeffizienten jedesmal genau dieselben Verhältnisse besitzen, daß also $\psi = \vartheta \varphi$ gilt, wo ϑ ein konstanter Faktor ist. (Die erstere Forderung wird z. B. gewiß erfüllt sein, wenn wir S und T derart auswählen, daß ihre Determinanten gleichen Wert haben und zudem in ihnen nach dieser Determinante als Modul je zwei entsprechende Koeffizienten immer gleichrestig sind.) Der Faktor ϑ wird als Quotient der Koeffizienten in ψ und φ wie diese Zahlen im Körper von α liegen. Setzen wir $(d-c)p_0 = p$, so gehen aus $T = PS$, $\psi = \vartheta \varphi$ vermöge $Q_j = Q_{j+p}$ ($j \geq j_0$) die Beziehungen

$$S_{j+p} = PS_j, \quad \varphi_{j+p} = \vartheta \varphi_j$$

für jeden Index $j \geq j_0$ hervor. Wir erhalten sodann allgemeiner

$$S_{j+f p} = P^f S_j, \quad \varphi_{j+f p} = \vartheta^f \varphi_j \quad (j \geq j_0)$$

für $f = 1, 2, 3, \dots$. Da in den Formen φ_j mit wachsendem Index j die Beträge der Koeffizienten jedenfalls nach Null abnehmen, muß $|\vartheta| < 1$ sein.

4. Wenn α komplex ist, bedeute α^0 die zu α konjugiert imaginäre Größe. Die $n-l$ Wurzeln der irreduziblen Gleichung für α außer α , bzw. außer α und α^0 mögen $\alpha', \alpha'', \dots, \alpha^{(n-l)}$ heißen. Ferner bezeichnen wir die zu einer Zahl ϑ oder einer Form ξ des Körpers von α konjugierten Zahlen oder Formen in den Körpern von $(\alpha^0), \alpha', \dots, \alpha^{(n-l)}$ analog durch

Hinzufügung oberer Indizes $(0), 1, \dots, n-l$. Durch die Substitution $P = S_{j_0+p} S_{j_0}^{-1}$ geht ξ in $\vartheta \xi$ und gehen daher wegen der Irreduzibilität der Gleichung n^{ten} Grades für α weiter $(\xi^0), \xi', \dots, \xi^{(n-l)}$ in $(\vartheta^0 \xi^0), \vartheta' \xi', \dots, \vartheta^{(n-l)} \xi^{(n-l)}$ über. Bedeutet nun t einen unbestimmten Parameter, E die identische Substitution, so gehen $\xi, \dots, \xi^{(n-l)}$ durch die Substitution $tE - P$ in $(t - \vartheta)\xi, \dots, (t - \vartheta^{(n-l)})\xi^{(n-l)}$ über und ist infolgedessen $|tE - P|$, d. h. die Determinante von $tE - P$, gleich dem Produkte $(t - \vartheta) \dots (t - \vartheta^{(n-l)})$. Diese in t identisch erfüllte Beziehung zeigt, daß ϑ der Gleichung $|tE - P| = 0$ genügt. Indem P eine ganzzahlige Substitution und ihre Determinante 1 ist, erweist sich dadurch ϑ als eine ganze Zahl und als eine Einheit im Körper von α .

5. Es sei nun a_0 eine solche ganze rationale Zahl, daß $a_0 \alpha$ eine ganze algebraische Zahl wird, so hat das Produkt $a_0^{n-1} \xi$ als Koeffizienten lauter ganze algebraische Zahlen und sind daher auch in jeder einzelnen Form $a_0^{n-1} \varphi_j$ die bezüglichlichen n Koeffizienten $a_0^{n-1} \varrho_k (k = 1, 2, \dots, n)$ stets lauter von Null verschiedene ganze algebraische Zahlen, also deren Normen im Körper von α stets dem Betrage nach ≥ 1 . Wegen der Eigenschaft a) der Substitutionen S_j liegt dabei jeder Betrag

$$\frac{|\varrho_k^{(h)}|}{r_j} \quad (h = 1, \dots, n-l; k = 1, 2, \dots, n)$$

nicht über einer gewissen von j unabhängigen Grenze. Verwenden wir nun die hierdurch gegebenen Ungleichungen für alle Indizes $h = 1, \dots, n-l$ mit Ausnahme eines beliebigen Index g dieser Reihe und berücksichtigen wir außerdem die Eigenschaft c) für S_j und, falls $l = 2$ ist, noch die Beziehung $|\varrho_k| = |\varrho_k^0|$, so gewinnen wir aus der Ungleichung

$$|Nm a_0^{n-1} \varrho_k| \geq 1$$

eine gewisse, von r_j unabhängige, positive untere Grenze für den einen darin übrig bleibenden Faktor $\frac{|\varrho_k^{(g)}|}{r_j}$. Danach befinden sich nun alle Beträge $\frac{|\varrho_k^{(g)}|}{r_j}$ in bezug auf die Form φ_j zwischen zwei bestimmten endlichen positiven, von r_j unabhängigen Grenzen, und werden infolgedessen weiter auch die Quotienten aus irgend zwei der je $n-l$ konjugierten Werte

$$\varrho_k', \varrho_k'', \dots, \varrho_k^{(n-l)} \quad (k = 1, 2, \dots, n)$$

bei sämtlichen Formen φ_j stets absolut genommen zwischen zwei, unabhängig von den Werten r_j feststehenden positiven endlichen Grenzen liegen. Beachten wir nun die Relationen $\varphi_{j+f} = \vartheta^f \varphi_j$ für $f = 1, 2, 3, \dots$, so zeigt sich schließlich, daß auch die Beträge der Quotienten aus irgend zwei der je $n-l$ Größen

$$\vartheta'^f, \vartheta''^f, \dots, (\vartheta^{(n-l)})^f$$

zwischen gewissen zwei festen positiven endlichen Grenzen liegen müssen, und zwar gelten diese Grenzen für alle Werte $f = 1, 2, 3, \dots$ auf einmal. Danach kann die Einheit ϑ nicht anders beschaffen sein, als daß für sie die Gleichungen

$$|\vartheta'| = |\vartheta''| = \dots = |\vartheta^{(n-1)}|$$

statthaben. Bezeichnen wir den gemeinsamen Wert dieser letzten Beträge mit η und setzen $|\vartheta| = \varepsilon$, womit im Falle $l = 2$ noch $|\vartheta^0|$ zusammenfällt, so geht die Gleichung $Nm\vartheta = 1$ in $\varepsilon^l \eta^{n-l} = 1$ über, und wegen $\varepsilon < 1$ folgt $\eta > 1$. Wir gelangen auf diese Weise zu dem Satze:

Damit eine algebraische Zahl n^{ten} Grades α eine periodische Substitutionenkette besitze, muß es im Körper von α eine Einheit ϑ von einem Betrage < 1 geben, für welche die konjugierten Zahlen in den konjugierten Körpern (abgesehen von der Zahl ϑ^0 in dem Körper der konjugiert imaginären Zahl α^0 , falls α komplex ist) sämtlich untereinander gleichen Betrag haben.

6. Die hier gefundene Bedingung ist zugleich hinreichend für das Vorhandensein einer periodischen Substitutionenkette zur Zahl α . Denn nehmen wir an, es existiere im Körper von α eine Einheit ϑ_0 von dem fraglichen Charakter. Es bedeute dann P_0 diejenige lineare Substitution, durch welche die n Formen $\xi, (\xi^0), \dots, \xi^{(n-1)}$ in die Formen $\vartheta_0 \xi, (\vartheta_0^0 \xi^0), \dots, \vartheta_0^{(n-1)} \xi^{(n-1)}$ übergehen; diese Substitution hat lauter rationale Koeffizienten und eine Determinante $= \pm 1$. Durch P_0^f , wenn f eine der Zahlen $1, 2, 3, \dots$ bedeutet, gehen dann $\xi, \dots, \xi^{(n-1)}$ in $\vartheta_0^f \xi, \dots, (\vartheta_0^{(n-1)})^f \xi^{(n-1)}$ über. Da diese Potenzen ϑ_0^f lauter ganze algebraische Zahlen sind, werden, wie leicht zu sehen ist, in allen jenen Substitutionen P_0^f die Koeffizienten solche ganze rationale Zahlen sein, daß ihre Nenner nicht über eine gewisse, durch die Größe α bestimmte, aber von den Exponenten f unabhängige Zahl hinausgehen, während zugleich ihre Determinanten durchweg $= \pm 1$ sind. Wir werden infolgedessen unter jenen unendlich vielen Substitutionen P_0^f gewiß irgend zwei solche, P_0^c und P_0^d ($d > c$), finden können, daß $P_0^d (P_0^c)^{-1} = P$ eine Substitution mit ganzzahligen Koeffizienten wird.

Setzen wir dann $\vartheta_0^{d-c} = \vartheta$, $|\vartheta|^{\frac{1}{n-1}} = \eta$, so haben wir in der Reihe

$$S_1 = E, \quad S_2 = P, \quad S_3 = P^2, \dots$$

eine periodische Substitutionenkette für die Zahl α mit den in 1. und 2. angegebenen Eigenschaften, wenn wir noch für die zugeordneten Größen r_j die Festsetzung $r_j = \eta^{j-1}$ ($j = 1, 2, 3, \dots$) treffen.

§ 2. Einheiten von besonderem Charakter.

7. Wir wollen jetzt die Forderung der Existenz der besonderen Einheit ϑ im Körper von α weiter verfolgen. Die ganze Funktion n^{ten} Grades in t :

$$F(t) = (t - \vartheta) \dots (t - \vartheta^{(n-1)})$$

hat *rationale ganze* Koeffizienten; unter ihren Wurzeln haben l den Betrag $\varepsilon < 1$ und $n - l$ den Betrag $\eta > 1$. Jeder im Bereiche der rationalen Zahlen irreduzible Faktor dieser Funktion $F(t)$ verschwindet für wenigstens eine der Zahlen $\vartheta, \dots, \vartheta^{(n-1)}$ und muß daher, wegen der Irreduzibilität der Gleichung mit den Wurzeln $\alpha, \dots, \alpha^{(n-1)}$, jedesmal für alle diese Zahlen $\vartheta, \dots, \vartheta^{(n-1)}$ verschwinden; infolgedessen ist $F(t)$ notwendig eine Potenz einer einzigen irreduziblen Funktion. Wegen der Beträge der Wurzeln sind nun offenbar nur diese beiden Fälle möglich: *Entweder* ist $F(t)$ selbst irreduzibel und bestimmt alsdann ϑ bereits den Körper von α , *oder* es ist α komplex, $l = 2$, aber $\vartheta = \vartheta^0$ reell und $F(t)$ das Quadrat einer irreduziblen Funktion; in letzterem Falle bestimmt ϑ einen reellen Unterkörper vom $\frac{n^{\text{ten}}}{2}$ Grade des komplexen Körpers von α . Wir bemerken noch, daß jede Potenz $\vartheta^2, \vartheta^3, \dots$ denselben Bedingungen genügt, wie sie hier für ϑ vorausgesetzt werden.

8. Nach einem Satze von Dirichlet gibt es in dem Körper der Zahl α , wenn nur $n > l$ ist, gewiß eine solche Einheit, deren Betrag < 1 ist. Daraus ersehen wir bereits, daß im Körper von α eine Einheit ϑ der hier verlangten Art sich gewiß in folgenden Fällen vorfindet:

- a) wenn α reell und $n = 2$ ist, b) wenn α reell ist, $n = 3$ und der Körper von α zwei komplexe konjugierte Körper besitzt,
c) wenn α komplex und $n = 3$ ist, d) wenn α komplex ist, $n = 4$ und der Körper von α lauter komplexe konjugierte Körper besitzt.

Denn in diesen Fällen besteht die Reihe $\vartheta', \dots, \vartheta^{(n-1)}$ entweder in einer einzigen reellen Zahl oder zwei konjugiert imaginären, im speziellen auch zwei gleichen reellen Zahlen. Weiter haben wir im Körper von α eine Einheit ϑ der verlangten Art jedenfalls auch in folgenden Fällen:

- e) wenn α komplex ist, $n = 4$ und der Körper von α einen reellen Unterkörper zweiten Grades hat, f) wenn α komplex ist, $n = 6$ und der Körper von α einen reellen Unterkörper dritten Grades besitzt, dessen zwei konjugierte Körper komplex sind.

Denn in diesen Fällen können wir für ϑ eine reelle Einheit von einem Betrage < 1 in dem betreffenden Unterkörper von α wählen, alsdann ist $\vartheta^0 = \vartheta$ und die Reihe $\vartheta', \dots, \vartheta^{(n-1)}$ besteht aus zwei gleichen reellen bzw. zwei gleichen Paaren konjugiert imaginärer Zahlen. Wir können jetzt den Satz beweisen:

Die hier aufgezählten sechs Fälle sind die einzigen, in denen der Körper von α eine Einheit ϑ der fraglichen Art aufweist, also die einzigen Fälle, in denen die Zahl α periodische Substitutionenketten besitzt.

9. Wir diskutieren zuerst den Fall einer reellen Zahl α ; hier ist $l = 1$, $\vartheta = \pm \varepsilon$, $\varepsilon \eta^{n-1} = 1$. Wir haben folgende Möglichkeiten ins Auge zu fassen:

a) Unter den Zahlen $\alpha', \dots, \alpha^{(n-1)}$ finden sich *wenigstens zwei reelle*, etwa $\alpha^{(h)}$ und $\alpha^{(k)}$. Dann sind auch $\vartheta^{(h)}$ und $\vartheta^{(k)}$ reell, und da diese Zahlen nicht einander gleich sein können, aber denselben Betrag haben, müßte $\vartheta^{(h)} = -\vartheta^{(k)}$ und daher $(\vartheta^{(h)})^2 = (\vartheta^{(k)})^2$ sein. Aber die Zahl $\vartheta^2 = \varepsilon^2$ ist von ihren $n - 1$ konjugierten Zahlen verschieden, sie genügt daher ebenfalls einer irreduziblen Gleichung n^{ten} Grades, und müßten daher ihre $n - 1$ konjugierten Zahlen auch untereinander durchweg verschieden sein. Danach ist dieser Fall unmöglich.

b) Unter den Zahlen $\alpha', \dots, \alpha^{(n-1)}$ kommt *nur eine reelle* Zahl, $\alpha^{(g)}$, vor. Für $n = 2$ liegt dann der oben unter 8. a) aufgeführte Fall vor. Ist $n > 2$, so haben wir unter jenen Zahlen weiter wenigstens ein Paar konjugiert imaginärer Zahlen, etwa $\alpha^{(h)}$ und $\alpha^{(k)}$. Die Zahl $\vartheta^2 = \varepsilon^2$ genügt einer irreduziblen Gleichung n^{ten} Grades; unter den Wurzeln dieser Gleichung ist weiter eine $= (\vartheta^{(g)})^2 = \eta^2$ und sind die $n - 2$ übrigen dem Betrage nach $= \eta^2$. Nun können wir eine Gleichung mit rationalen Koeffizienten vom Grade $\frac{n(n-1)}{2}$ angeben, welche die Produkte aus je zwei verschiedenen der n Größen $\vartheta, \vartheta', \dots, \vartheta^{(n-1)}$ zu Wurzeln hat. Diese Gleichung besitzt $n - 1$ Wurzeln vom Betrage $\varepsilon \eta$, die übrigen Wurzeln vom Betrage η^2 , darunter insbesondere die Wurzel $\vartheta^{(h)} \vartheta^{(k)} = \eta^2$, sie müßte also auch alle anderen Wurzeln jener irreduziblen Gleichung n^{ten} Grades für η^2 besitzen; sie hätte aber, da $\varepsilon < 1 < \eta$ ist, gewiß nicht die Wurzel ε^2 . Danach ist dieser Fall für $n > 2$ unmöglich.

c) Die Zahlen $\alpha', \dots, \alpha^{(n-1)}$ sind *sämtlich komplex*, sie zerfallen dann in $\frac{n-1}{2}$ Paare konjugiert imaginärer Größen. Für $n = 3$ liegt der oben unter 8. b) aufgeführte Fall vor. Jetzt sei $n > 3$. Wir bilden die Gleichung $\frac{n(n-1)}{2}^{\text{ten}}$ Grades mit rationalen Koeffizienten, welche als Wurzeln die Produkte aus je zwei der n Größen $\vartheta^{-n+1}, \vartheta'^{-n+1}, \dots, (\vartheta^{(n-1)})^{-n+1}$ hat. Diese Gleichung besitzt $n - 1$ Wurzeln vom Betrage $(\varepsilon \eta)^{-n+1} = \eta^{(n-1)(n-2)}$ und im übrigen lauter Wurzeln vom Betrage $\eta^{-2(n-1)} = \varepsilon^2$, darunter $\frac{n-1}{2}$ Wurzeln $= \varepsilon^2 = \vartheta^2$; sie müßte daher auch alle die Größen $\vartheta'^2, \dots, (\vartheta^{(n-1)})^2$ vom Betrage η^2 zu Wurzeln besitzen, es müßte also $\eta^2 = \eta^{(n-1)(n-2)}$, d. h. $n = 3$ sein. Für $n > 3$ ist danach dieser Fall unmöglich.

10. Wir behandeln jetzt weiter den *Fall einer komplexen Zahl α* ; hier ist $l = 2$, $\varepsilon^2 \eta^{n-2} = 1$.

Machen wir zunächst die Annahme, daß $\vartheta = \vartheta^0$, also reell ist. Die Größe ϑ ist dann Wurzel einer irreduziblen Gleichung $\frac{n}{2}^{\text{ten}}$ Grades. Der Körper von α besitzt also einen reellen Unterkörper vom Grade $\frac{n}{2}$, und in diesem soll ϑ eine Einheit von einem Betrage < 1 sein, für welche die konjugierten Zahlen in den konjugierten Körpern sämtlich untereinander gleiche Beträge haben. Wir können daher die in 9. gemachten Ausführungen verwenden, und es muß entweder $\frac{n}{2} = 2$ sein oder aber $\frac{n}{2} = 3$ und dabei der Körper von ϑ zwei komplexe konjugierte Körper aufweisen. Wir kommen damit auf die oben unter 8. e) und 8. f) aufgezählten Umstände für den Körper von α .

Wir nehmen jetzt andererseits an, daß $\vartheta \neq \vartheta^0$ sei. Alsdann genügt ϑ einer irreduziblen Gleichung n^{ten} Grades und bestimmt bereits völlig den Körper von α . Wir unterscheiden wieder drei Fälle:

a) Unter den Zahlen $\alpha', \dots, \alpha^{(n-2)}$ sind *wenigstens zwei reelle* vorhanden, $\alpha^{(h)}$ und $\alpha^{(k)}$. Dann sind auch $\vartheta^{(h)}$ und $\vartheta^{(k)}$ reell, und da sie gleichen Betrag haben, aber verschieden sind, kann nur $\vartheta^{(h)} = -\vartheta^{(k)}$ sein. Alsdann ist $(\vartheta^{(h)})^2 = (\vartheta^{(k)})^2$. Die rationale Gleichung n^{ten} Grades mit den Wurzeln $\vartheta^2, \dots, (\vartheta^{(n-2)})^2$ hat daher lauter Doppelwurzeln und muß infolgedessen $\vartheta^2 = (\vartheta^0)^2$ sein. Die Potenz ϑ^2 bestimmt somit einen reellen Unterkörper $\frac{n}{2}^{\text{ten}}$ Grades für den Körper von α , und da überdies $(\vartheta^{(h)})^2$ reell ist, kann nach dem vorhin Bemerkten hier nur der unter 8. e) aufgeführte Fall mit $n = 4$ vorliegen.

b) Unter den Zahlen $\alpha', \dots, \alpha^{(n-2)}$ ist *nur eine reelle* Zahl, $\alpha^{(s)}$, vorhanden. Für $n = 3$ liegt der unter 8. c) genannte Fall vor. Ist $n > 3$, so haben wir unter jenen Zahlen noch $\frac{n-3}{2}$ Paare konjugiert imaginärer Größen. Es ist $\vartheta^{(s)} = \pm \eta$; da $-\vartheta^{(s)}$ hier nicht derselben Gleichung mit rationalen Koeffizienten wie $\vartheta^{(s)}$ genügt, muß notwendig auch $\vartheta^0 \neq -\vartheta$, also $(\vartheta^0)^2 \neq \vartheta^2$ und daher $\vartheta^2 \neq \pm \varepsilon^2$, $(\vartheta^0)^2 \neq \pm \varepsilon^2$ sein. Danach ist die rationale Gleichung n^{ten} Grades mit den Wurzeln $\vartheta^2, \dots, (\vartheta^{(n-2)})^2$ irreduzibel und unter den Wurzeln dieser Gleichung sind zwei nicht reelle Wurzeln vom Betrage ε^2 und ist ferner eine Wurzel $= \eta^2$ vorhanden. Bilden wir nun die rationale Gleichung $\frac{n(n-1)}{2}^{\text{ten}}$ Grades, welche die Produkte aus je zwei der n Größen $\vartheta, \dots, \vartheta^{(n-2)}$ zu Wurzeln hat, so besitzt diese Gleichung eine Wurzel $= \varepsilon^2$, sodann $2(n-2)$ Wurzeln vom Betrage $\varepsilon\eta$, die übrigen Wurzeln vom Betrage η^2 und darunter $\frac{n-3}{2}$ Wurzeln $= \eta^2$. Wegen der

letzteren Wurzeln müßte sie aber alle Wurzeln jener irreduziblen Gleichung für η^2 besitzen. Danach ist dieser Fall für $n > 3$ unmöglich.

c) Die Zahlen $\alpha', \dots, \alpha^{(n-2)}$ sind sämtlich komplex, zerfallen also in $\frac{n-2}{2}$ Paare konjugiert imaginärer Größen; n ist hier gerade. Für $n = 4$ liegt der oben unter 8. d) aufgeführte Fall vor. Jetzt sei $n \geq 6$. Bilden wir die Gleichung $\frac{n(n-1)}{2}$ Grades, welche die Produkte aus je zwei der n Größen $\vartheta, \dots, \vartheta^{(n-2)}$ zu Wurzeln hat, so besitzt diese Gleichung mit rationalen Koeffizienten eine Wurzel $= \varepsilon^2$, sodann $2(n-2)$ Wurzeln vom Betrage $\varepsilon\eta$, die übrigen Wurzeln vom Betrage η^2 , darunter $\frac{n-2}{2}$ gleich η^2 . Bilden wir andererseits die Gleichung $\frac{n(n-1)}{2}$ Grades, deren Wurzeln die $-\frac{n-2}{2}$ ten Potenzen der Wurzeln dieser letzten Gleichung sind, so hat diese neue Gleichung mit rationalen Koeffizienten eine Wurzel $= \varepsilon^{-(n-2)} = \eta^{\frac{(n-2)^2}{2}}$, sodann $2(n-2)$ Wurzeln vom Betrage $(\varepsilon\eta)^{\frac{-(n-2)}{2}} = \eta^{\frac{(n-2)(n-4)}{4}}$, die übrigen Wurzeln vom Betrage $\eta^{-(n-2)} = \varepsilon^2$, darunter $\frac{n-2}{2}$ gleich ε^2 . Diese zweite Gleichung besitzt nun keine Wurzel vom Betrage $\varepsilon\eta$, und wenn wir uns zuerst $n > 6$ denken, auch keine Wurzel vom Betrage η^2 . Im Falle $n > 6$ könnte daher der gemeinsame Faktor der beiden eben gebildeten Gleichungen nur die eine Wurzel ε^2 besitzen, es müßte dann also $\varepsilon^2 = \vartheta\vartheta^0$ rational sein; nun wäre aber ε^2 ebenso wie ϑ eine Einheit, eine algebraische Zahl von der Norm ± 1 , es müßte daher notwendig $\varepsilon^2 = 1$ sein, was gegen die Voraussetzung $\varepsilon < 1$ wäre. Also ist die Annahme $n > 6$ hier unzulässig.

Im Falle $n = 6$ endlich hat die zuerst erwähnte Gleichung eine Wurzel $= \varepsilon^2$, 8 Wurzeln vom Betrage $\varepsilon\eta$, 6 vom Betrage η^2 , die an zweiter Stelle gebildete Gleichung eine Wurzel $= \eta^8$, 8 Wurzeln vom Betrage η^2 , 6 vom Betrage ε^2 , darunter 2 gleich ε^2 . Die im Bereiche der rationalen Zahlen irreduzible Gleichung mit ε^2 als Wurzel kann danach, da sie in diesen beiden Gleichungen als Faktor eingeht, außer ε^2 nur Wurzeln vom Betrage η^2 enthalten, und sie wird wegen $\varepsilon^2\eta^4 = 1$ und da $\varepsilon^2 = \vartheta\vartheta^0$ jedenfalls eine Einheit vorstellt, im ganzen zwei solcher Wurzeln enthalten; diese zwei Wurzeln können dann einander weder gleich noch entgegengesetzt, also auch nicht reell $= \pm \eta^2$ sein, sie müssen komplex und zueinander konjugiert imaginär sein. Nehmen wir an, daß hier ϑ' mit ϑ'' und ϑ''' mit $\vartheta^{(4)}$ konjugiert imaginär sind, so können wir annehmen, indem wir noch die Bezeichnungen von ϑ''' und $\vartheta^{(4)}$ vertauschen dürfen, die Wurzeln jener Gleichung für ε^2 seien $\vartheta\vartheta^0$, $\vartheta'\vartheta'''$, $\vartheta''\vartheta^{(4)}$.

Die Größe ε^2 bestimmt danach einen kubischen Körper, dessen zwei konjugierte Körper komplex sind. Denken wir uns jetzt die im Bereiche der rationalen Zahlen irreduzible ganze Funktion $F(t)$, welche für $t = \vartheta$ verschwindet, im Körper von ε^2 in irreduzible Faktoren zerlegt, und es sei $G(t)$ derjenige Faktor darunter, welcher die Wurzel $t = \vartheta$ erhält. Da ε^2 reell ist, bekommt $G(t)$ ebenfalls lauter reelle Koeffizienten und wird daher mit der Wurzel ϑ auch die konjugiert imaginäre Größe $\vartheta^0 = \frac{\varepsilon^2}{\vartheta}$ als Wurzel besitzen, so daß auch $G\left(\frac{\varepsilon^2}{\vartheta}\right) = 0$ ist. Alsdann muß die Gleichung $G\left(\frac{\varepsilon^2}{t}\right) = 0$ überhaupt für jede Wurzel der im Körper von ε^2 irreduziblen Gleichung $G(t) = 0$ bestehen. Die Größen $\frac{\varepsilon^2}{\vartheta'}$, $\frac{\varepsilon^2}{\vartheta''}$, $\frac{\varepsilon^2}{\vartheta'''}$, $\frac{\varepsilon^2}{\vartheta^{(4)}}$ aber besitzen sämtlich den Betrag $\frac{\varepsilon^2}{\eta} \neq \eta$ und $\neq \varepsilon$ und sind daher nicht Wurzeln von $F(t) = 0$, also auch nicht Wurzeln von $G(t) = 0$, daher kann $G(t)$ auch keine der Größen ϑ' , ϑ'' , ϑ''' , $\vartheta^{(4)}$ als Wurzel haben; somit können wir einfach $G(t) = (t - \vartheta)(t - \vartheta^0)$ schreiben. Danach ist ϑ Wurzel einer quadratischen Gleichung im Körper von ε^2 und besitzt der Körper sechsten Grades von ϑ , d. i. der Körper von α , in dem Körper von ε^2 einen reellen Unterkörper dritten Grades, dessen zwei konjugierte Körper komplex sind. Wir kommen also auf den oben unter 8. f) aufgezählten Fall.

Wir können die Bildungsweise des Körpers von α unter den hier angenommenen Umständen noch genauer festlegen. Die zu $G(t)$ konjugierten Funktionen in den Körpern von $\vartheta'\vartheta'''$ und $\vartheta''\vartheta^{(4)}$ werden $(t - \vartheta')(t - \vartheta''')$ und $(t - \vartheta'')(t - \vartheta^{(4)})$ sein. Da $|\vartheta'| = |\vartheta'''|$ ist, wird $\frac{\vartheta' - \vartheta'''}{\vartheta' + \vartheta'''}$ rein imaginär, also $\left(\frac{\vartheta' - \vartheta'''}{\vartheta' + \vartheta'''}\right)^2$ eine negative reelle Größe sein. Diese Größe liegt wie $\vartheta' + \vartheta'''$ und $\vartheta'\vartheta'''$ in dem Körper von $\vartheta'\vartheta'''$; da sie nun reell ist, wird sie identisch mit ihrer konjugierten Größe in dem konjugiert imaginären Körper von $\vartheta''\vartheta^{(4)}$ und muß daher rational sein und daher gleichzeitig auch gleich der konjugierten Größe $\left(\frac{\vartheta - \vartheta^0}{\vartheta + \vartheta^0}\right)^2$ im Körper von $\vartheta\vartheta^0$. Danach ist $\frac{\vartheta}{\vartheta^0}$ entweder $= \frac{\vartheta'}{\vartheta'''}$ oder $= \frac{\vartheta''}{\vartheta^{(4)}}$. Da wir die Bezeichnung der Paare ϑ' , ϑ'' und ϑ''' , $\vartheta^{(4)}$ vertauschen dürfen, können wir annehmen, es sei $\frac{\vartheta}{\vartheta^0} = \frac{\vartheta'}{\vartheta''}$; letzterer Wert ist weiter $= \frac{\vartheta^{(4)}}{\vartheta''}$. Setzen wir $\frac{\vartheta}{\vartheta^0} = \delta$, so ist $\delta = \frac{\vartheta^2}{\vartheta\vartheta^0}$ und sind die konjugierten Zahlen dazu $\frac{\vartheta'^2}{\vartheta'\vartheta''} = \delta$, $\frac{(\vartheta^{(4)})^2}{\vartheta''\vartheta^{(4)}} = \delta$ und die drei hierzu reziproken Werte $= \frac{1}{\delta}$. Dabei hat δ den Betrag 1 und ist wie ϑ eine Einheit; danach ist entweder $\delta = -1$, oder es ist δ eine solche Einheits-

wurzel, die einen Körper zweiten Grades bestimmt. Im ersteren Falle ist $\vartheta^0 = -\vartheta$, $\vartheta = \pm i\varepsilon$, $\vartheta^2 = (\vartheta^0)^2$. Im zweiten Falle kämen für δ zunächst die dritten, vierten, sechsten Einheitswurzeln in Betracht. Nun folgt $\vartheta = \delta^{\frac{1}{2}}\varepsilon$, $\vartheta^0 = \frac{1}{\delta^{\frac{1}{2}}}\varepsilon$ und weiter unter Verwendung von $\varepsilon\eta^2 = \varepsilon\vartheta'\vartheta'' = 1$

die Beziehung

$$Nm(\vartheta + \vartheta^0) = (\vartheta + \vartheta^0)(\vartheta' + \vartheta'')(\vartheta'' + \vartheta^{(4)}) = \frac{\delta+1}{\delta^{\frac{1}{2}}}\left(1 + \frac{1}{\delta}\right)(1 + \delta).$$

Danach muß noch $\frac{\delta+1}{\delta^{\frac{1}{2}}}$ rational sein, und solches trifft nur zu, wenn δ

eine dritte Einheitswurzel vorstellt, $\delta^3 = 1$ ist. Dann folgt endlich $\vartheta = \pm \frac{\varepsilon}{\delta}$, $\vartheta^0 = \pm \delta\varepsilon$, $\vartheta^3 = (\vartheta^0)^3$ und $\vartheta + \vartheta^0 = \mp \varepsilon$, so daß der Körper von ε^2 auch die Größe ε selbst enthält, und der Körper von α aus dem Körper dritten Grades von ε und dem Körper zweiten Grades von $\delta = \frac{-1 \pm \sqrt{-3}}{2}$ zusammengesetzt ist.

§ 3. Die komplexen kubischen Irrationalzahlen.

11. In den Fällen, wo für die algebraische Zahl α periodische Substitutionenketten möglich sind, entsteht nun die Aufgabe, eine solche Kette für α bereits herzustellen, wenn allein α seinem Werte nach gegeben ist, die konjugierten algebraischen Zahlen von α indes noch unbekannt sind. Bei den reellen algebraischen Zahlen zweiten Grades wird gerade durch die periodische Entwicklung in einen gewöhnlichen Kettenbruch das hier Verlangte geleistet. Wir wollen nun zeigen, daß auch noch in einem anderen Falle, nämlich, wenn es sich um eine komplexe Größe α handelt, welche Wurzel einer irreduziblen Gleichung dritten Grades sein soll, der hier gestellten Forderung entsprochen werden kann, und wir kommen dadurch zu einem völlig analogen Kriterium für die komplexen algebraischen Zahlen dritten Grades, wie es durch Lagrange für die reellen algebraischen Zahlen zweiten Grades in der Periodizität der Kettenbruchentwicklung nachgewiesen worden ist.

Es sei jetzt α eine komplexe Größe, welche einer im Bereiche der rationalen Zahlen irreduziblen Gleichung dritten Grades genügt, so ist mit α ohne weiteres auch die konjugiert imaginäre Größe α^0 gegeben; dagegen haben wir uns der Kenntnis der dritten reellen Wurzel α' jener Gleichung vorläufig zu entschlagen. Wir setzen

$$\xi = x_1 + \alpha x_2 + \alpha^2 x_3.$$

Zu jeder ganzen rationalen Zahl $r \geq 1$ bestimmen wir eine Substitution S :

$$x_h = s_h^{(1)}y_1 + s_h^{(2)}y_2 + s_h^{(3)}y_3 \quad (h = 1, 2, 3),$$

wobei jeder der Koeffizienten $s_h^{(k)}$ aus den Zahlen $0, \pm 1, \pm 2, \dots, \pm r$ entnommen und die Determinante $\neq 0$ sein soll, derart, daß dabei in der Form

$$\varphi = \varrho_1 y_1 + \varrho_2 y_2 + \varrho_3 y_3,$$

in welche ξ durch S übergeht, zunächst $|\varrho_1|$ möglichst klein, nächst dem $|\varrho_2|$ möglichst klein, endlich $|\varrho_3|$ möglichst klein ausfällt. Bei den einzelnen Systemen $s_1^{(k)}, s_2^{(k)}, s_3^{(k)}$ haben wir immer die Wahl unter Paaren entgegengesetzter Systeme x_1, x_2, x_3 und $-x_1, -x_2, -x_3$, und wir wollen die zusätzliche Annahme machen, daß in jeder Vertikalreihe $s_1^{(k)}, s_2^{(k)}, s_3^{(k)}$ von S die letzte von Null verschiedene Zahl stets > 0 ist. *Alsdann ist die betreffende Substitution S durch die Zahl r vollkommen eindeutig bestimmt.* Wir können nämlich niemals für zwei Systeme x_1, x_2, x_3 , die $\neq 0, 0, 0$, voneinander verschieden und auch nicht einander entgegengesetzt sind, $\xi = \varrho$, $\xi = \sigma$ und dabei $|\varrho| = |\sigma|$ finden. Denn alsdann wäre $\frac{\varrho \varrho^0}{\sigma \sigma^0} = 1$ und würde die Betrachtung der Norm von $\frac{\varrho}{\sigma}$ in bezug auf den Körper von α dazu führen, daß die zu $\frac{\varrho}{\sigma}$ konjugierte Zahl im Körper von α' rational wäre. Dann aber wäre auch $\frac{\varrho}{\sigma}$ selbst rational, also notwendig $= \pm 1$, und müßten die vorausgesetzten zwei Systeme x_1, x_2, x_3 eben entweder gleich oder entgegengesetzt sein. Die in dieser Weise durch r völlig bestimmte Substitution S entspricht, wie bereits in 1. erwähnt wurde, den dort aufgezählten Umständen; wir wollen von dieser Substitution S sagen, sie *gehöre* zur Zahl r .

Wir setzen jetzt $r_1 = 1$ und ermitteln die zu r_1 gehörende Substitution S_1 , die Substitution kann auch noch zu $r = 2, 3, \dots$ gehören, es sei $r_2 - 1$ die größte ganze Zahl, zu der sie gehört; sodann sei S_2 die zu r_2 gehörende Substitution, $r_3 - 1$ die größte ganze Zahl, zu der S_2 gehört; S_3 die zu r_3 gehörende Substitution usw. Alsdann gilt der folgende Satz:

Die in dieser Art hergestellte Substitutionskette $S_1, S_2, S_3, \dots$ für die komplexe kubische Irrationalzahl α ist stets periodisch.

12. In der Tat, im Körper von α können wir stets eine Einheit ϑ_0 angeben, deren Betrag < 1 ist, und noch so, daß $\vartheta_0' > 0$ ist; alsdann können wir, wie aus den Betrachtungen in 6. hervorgeht, eine solche Potenz dieser Einheit $\vartheta_0^f = \vartheta(f > 0)$ ermitteln, daß die lineare Substitution P , durch welche die Formen ξ, ξ^0, ξ' in $\vartheta \xi, \vartheta^0 \xi^0, \vartheta' \xi'$ übergehen, lauter *ganzzahlige* Koeffizienten hat; dabei ist $\vartheta \vartheta^0 \vartheta' > 0$ und die Determinante von P gleich 1.

Wir schreiben nun die Auflösung von

$$\xi = x_1 + \alpha x_2 + \alpha^2 x_3, \quad \xi^0 = x_1 + \alpha^0 x_2 + (\alpha^0)^2 x_3, \quad \xi' = x_1 + \alpha' x_2 + \alpha'^2 x_3$$

in der Form

$$x_h = \beta_h \xi + \beta_h^0 \xi^0 + \beta_h' \xi' \quad (h=1, 2, 3);$$

dabei sind β_h , β_h^0 , β_h' jedesmal konjugierte Zahlen in den Körpern von α , α^0 , α' .

Es sei jetzt S :

$$x_h = s_h^{(1)} y_1 + s_h^{(2)} y_2 + s_h^{(3)} y_3 \quad (h=1, 2, 3)$$

eine Substitution der oben gebildeten Kette und r die niedrigste Zahl, zu der diese Substitution S gehört, also r der größte Wert unter den Beträgen der 9 Koeffizienten $s_h^{(k)}$, und es gehe ξ durch S in

$$\varphi = \varrho_1 y_1 + \varrho_2 y_2 + \varrho_3 y_3$$

über, so folgt

$$s_h^{(1)} y_1 + s_h^{(2)} y_2 + s_h^{(3)} y_3 = \beta_h \varphi + \beta_h^0 \varphi^0 + \beta_h' \varphi'$$

und daraus

$$s_h^{(k)} = \beta_h \varrho_k + \beta_h^0 \varrho_k^0 + \beta_h' \varrho_k' \quad (h, k=1, 2, 3).$$

Nach der in 1. erwähnten Eigenschaft c) können wir eine, von α allein abhängende, von r aber völlig unabhängige positive Konstante M angeben, so daß $|\varrho_1| r^{\frac{1}{2}}$, $|\varrho_2| r^{\frac{1}{2}}$, $|\varrho_3| r^{\frac{1}{2}}$ bei jedem Werte von r stets $\leq M$ sind, und nach 5. gelangen wir dann durch Betrachtung der Normen von ϱ_1 , ϱ_2 , ϱ_3 unter Berücksichtigung des Umstandes, daß die Norm einer von Null verschiedenen *ganzen* Zahl stets ≥ 1 ist, zu einer weiteren, von r unabhängigen positiven Konstante, die wir $\frac{M}{N}$ schreiben wollen, so daß $|\varrho_1'| r^{-1}$, $|\varrho_2'| r^{-1}$, $|\varrho_3'| r^{-1}$ stets $\geq \frac{M}{N}$ sind.

Auf diese Ungleichungen gestützt, können wir jetzt gewisse Eigenschaften für die Substitution $T = PS$ nachweisen, sowie nur r eine bestimmte Größe übersteigt. Durch $T = PS$ gehen die Formen ξ , ξ^0 , ξ' in $\vartheta \varphi$, $\vartheta^0 \varphi^0$, $\vartheta' \varphi'$ über; ist

$$x_h = t_h^{(1)} y_1 + t_h^{(2)} y_2 + t_h^{(3)} y_3 \quad (h=1, 2, 3)$$

der Ausdruck dieser Substitution T , so folgt daher

$$t_h^{(k)} = \vartheta \beta_h \varrho_k + \vartheta^0 \beta_h^0 \varrho_k^0 + \vartheta' \beta_h' \varrho_k' \quad (h, k=1, 2, 3).$$

Alsdann haben wir

$$\frac{t_h^{(k)}}{s_h^{(k)}} = \vartheta' \frac{1 + \frac{\vartheta \beta_h}{\vartheta' \beta_h'} \frac{\varrho_k}{\varrho_k'} + \frac{\vartheta^0 \beta_h^0}{\vartheta' \beta_h'} \frac{\varrho_k^0}{\varrho_k'^0}}{1 + \frac{\beta_h}{\beta_h'} \frac{\varrho_k}{\varrho_k'} + \frac{\beta_h^0}{\beta_h'} \frac{\varrho_k^0}{\varrho_k'^0}}.$$

Setzen wir $|\vartheta| = \varepsilon$, wobei $\vartheta' = \varepsilon^{-2}$ wird, und bezeichnen wir noch den größten Wert unter den Beträgen $\left| \frac{\beta_h}{\beta_h'} \right|$ für $h=1, 2, 3$ mit γ , so geht

hieraus auf Grund der erwähnten Ungleichungen, *wofern* $2\gamma N r^{-\frac{3}{2}} < 1$ ist, einerseits

$$\frac{t_h^{(k)}}{s_h^{(k)}} \geq \vartheta' \frac{1 - 2\varepsilon^3 \gamma N r^{-\frac{3}{2}}}{1 + 2\gamma N r^{-\frac{3}{2}}} > \vartheta' \left(1 - 2(\varepsilon^3 + 1) \gamma N r^{-\frac{3}{2}} \right),$$

andererseits

$$\frac{t_h^{(k)}}{s_h^{(k)}} \leq \vartheta' \frac{1 + 2\varepsilon^3 \gamma N r^{-\frac{3}{2}}}{1 - 2\gamma N r^{-\frac{3}{2}}} < \vartheta' \left(1 + 2(\varepsilon^3 + 1) \gamma N r^{-\frac{3}{2}} \right)$$

hervor. Wir nehmen nunmehr die Zahl r überhaupt so groß an, daß die stärkere Bedingung

$$(I) \quad 2\vartheta'(\varepsilon^3 + 1) \gamma N r^{-\frac{1}{2}} < \frac{1}{2}$$

erfüllt ist. Alsdann finden wir, indem alle Zahlen $s_h^{(k)}$ dem Betrage nach $\leq r$ und wenigstens eine darunter $= \pm r$ ist, die Beträge der Zahlen $t_h^{(k)}$ sämtlich $< \vartheta' r + \frac{1}{2}$ und wenigstens einen unter diesen Beträgen $> \vartheta' r - \frac{1}{2}$; es wird danach der größte unter diesen Beträgen der 9 Koeffizienten $t_h^{(k)}$ gleich der an $\vartheta' r$ nächstgelegenen ganzen Zahl sein, die wir mit $\bar{r}$ bezeichnen wollen. Ferner finden wir alle Zahlen $s_h^{(k)} \neq 0$ und die Quotienten $\frac{t_h^{(k)}}{s_h^{(k)}} > 0$, und wird daher gewiß in T in jedem Systeme $t_1^{(k)}, t_2^{(k)}, t_3^{(k)}$ die letzte von Null verschiedene Zahl, nämlich $t_3^{(k)} > 0$ sein, wie wir das entsprechende für S voraussetzen.

13. Aus den nämlichen Relationen, die wir soeben behandelten, ersehen wir andererseits die folgende Tatsache: Es sei T eine Substitution der in 11. für die Zahl α gebildeten Kette, $\bar{r}$ die kleinste Zahl, zu der T gehört, und $\bar{r} > \vartheta'$. Bilden wir die Substitution $S = P^{-1}T$, so bestehen zwischen je zwei entsprechenden Koeffizienten $t_h^{(k)}$ von T und $s_h^{(k)}$ von S , sowie nur $\bar{r}$ der Ungleichung $2\gamma N \bar{r}^{-\frac{3}{2}} < 1$ genügt, stets die Beziehungen

$$\frac{1}{\vartheta'} \left(1 + 2 \left(\frac{1}{\varepsilon^3} + 1 \right) \gamma N \bar{r}^{-\frac{3}{2}} \right) > \frac{s_h^{(k)}}{t_h^{(k)}} > \frac{1}{\vartheta'} \left(1 - 2 \left(\frac{1}{\varepsilon^3} + 1 \right) \gamma N \bar{r}^{-\frac{3}{2}} \right).$$

Setzen wir nun für $\bar{r}$ die stärkere Ungleichung

$$(II) \quad \frac{2}{\vartheta'} \left(\frac{1}{\varepsilon^3} + 1 \right) \gamma N \bar{r}^{-\frac{1}{2}} < \frac{1}{2}$$

voraus, so wird daher der größte unter den Beträgen der Koeffizienten $s_h^{(k)}$ mit der an $\vartheta'^{-1}\bar{r}$ nächstgelegenen ganzen Zahl übereinstimmen.

14. Jetzt denken wir uns wieder, wie in 12., S als irgendeine Substitution jener Kette, und es soll dabei sowohl die niedrigste Zahl r , zu der S gehört, der Bedingung (I) entsprechen, wie auch die an $\vartheta' r$ nächstgelegene ganze Zahl $\bar{r}$ die Bedingung (II) erfüllen. Alsdann können wir behaupten, daß $T = PS$ jedesmal ebenfalls eine Substitution in jener Kette ist und zu $\bar{r}$ als niedrigster Zahl gehört. Denn in T sind wegen (I) gewiß alle Koeffizienten $t_h^{(k)}$ dem Betrage nach $\leq \bar{r}$, und zudem ist in jeder Vertikalreihe von T die letzte von Null verschiedene Zahl > 0 . Wäre nun die zur Zahl $\bar{r}$ gehörende Substitution der Kette, T^* , von T ver-

schieden und würde durch sie ξ in $\vartheta(\varrho_1^*y_1 + \varrho_2^*y_2 + \varrho_3^*y_3)$ übergehen, so müßte dabei entweder $|\varrho_1^*| < |\varrho_1|$ oder $|\varrho_1^*| = |\varrho_1|$, $|\varrho_2^*| < |\varrho_2|$ oder $|\varrho_1^*| = |\varrho_1|$, $|\varrho_2^*| = |\varrho_2|$, $|\varrho_3^*| < |\varrho_3|$ sein. In der Substitution $S^* = P^{-1}T^*$ würden dann wegen (II) alle Koeffizienten dem Betrage nach

$$< \frac{\bar{r}}{\vartheta'} + \frac{1}{2} < \frac{1}{\vartheta'} \left(\vartheta' r + \frac{1}{2} \right) + \frac{1}{2} < r + 1,$$

also als ganze Zahlen auch $\leq r$ sein, und würde ξ durch S in $\varrho_1^*y_1 + \varrho_2^*y_2 + \varrho_3^*y_3$ übergehen, wobei die soeben erwähnten Bedingungen statthätten. Danach könnte S nicht die zur Zahl r gehörende Substitution der Kette sein.

Ist andererseits T eine solche Substitution jener Kette, die zu einer Zahl $\bar{r} > \vartheta'$ als niedrigster Zahl gehört, und erfüllt sowohl $\bar{r}$ die Bedingung (II) wie die an $\vartheta'^{-1}\bar{r}$ nächstgelegene ganze Zahl r die Bedingung (I), so erkennen wir ganz analog, daß auch $S = P^{-1}T$ jedesmal eine Substitution der Kette ist und zu r als niedrigster Zahl gehört.

15. Wir wählen jetzt in der Reihe $r_1, r_2, r_3, \dots$ die Größe r_{j_0} derart aus, daß die Ungleichung (I) mit $r = r_{j_0}$ und die Ungleichung (II) mit $\bar{r} = \vartheta' r_{j_0} - \frac{1}{2}$ erfüllt ist. Alsdann kommt mit der Substitution S_{j_0} in der Kette an irgendeiner späteren Stelle die Substitution PS_{j_0} vor, sie sei etwa $= S_{j_0+p}$; dabei wird r_{j_0+p} die an $\vartheta' r_{j_0}$ nächstgelegene ganze Zahl. Jetzt gilt (I) auch für jede Größe $r = r_j$, wobei $j > j_0$ ist, und (II) auch für jede Größe $\bar{r} = r_{j+p}$, wobei $j > j_0$ ist. Mit jeder Substitution S_j ($j > j_0$) wird daher auch die Substitution PS_j der Kette angehören und zwar als ein um so späteres Glied, je größer j ist, und andererseits wird mit jeder Substitution S_{j+p} ($j > j_0$) auch $P^{-1}S_{j+p}$ der Kette angehören, und zwar als ein um so späteres Glied, je größer j ist. Daraus folgt dann, daß die Substitutionen $PS_{j_0}, PS_{j_0+1}, PS_{j_0+2}, \dots$ sämtlich in der Kette auftreten, und zwar jede Substitution *später* als die hier vorher genannte, daß aber andererseits in der Kette keine Substitution *zwischen* diesen einzelnen Substitutionen vorkommen, kann, daß mithin allgemein

$$PS_j = S_{j+p}$$

für $j = j_0, j_0 + 1, j_0 + 2, \dots$ ist. Aus diesen Beziehungen entnehmen wir

$$S_j^{-1}S_{j+1} = S_{j+p}^{-1}S_{j+p+1}$$

für $j \geq j_0$, d. h. die Kette $S_1, S_2, S_3, \dots$ ist in der Tat periodisch, w. z. z. w.

Es kann bewiesen werden, daß für eine jede Substitution S der Kette die Determinante nur die Werte ± 1 oder ± 2 haben kann, und auf Grund dieser Tatsache läßt sich ein einfacher Algorithmus zur sukzessiven Bildung der Substitutionen der Kette aufstellen, wie ich an einer anderen Stelle auseinandersetzen werde.

Zürich, 1902.



Verlag B. G. Teubner, Leipzig.

Bleichner & Leykauf, Wien, hel. & imp.

H. Minkowski

NACH EINER AUFNAHME AUS DER ZEIT DER PARISER PREISARBEIT

GESAMMELTE ABHANDLUNGEN

VON

HERMANN MINKOWSKI

UNTER MITWIRKUNG VON

ANDREAS SPEISER UND **HERMANN WEYL**

HERAUSGEGEBEN VON

DAVID HILBERT

ZWEITER BAND

INHALT DES ZWEITEN BANDES.

Zur Geometrie der Zahlen (Fortsetzung).

Seite

- XIX. Dichteste gitterförmige Lagerung kongruenter Körper. 3
 (Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen,
 mathematisch-physikalische Klasse, 1904, S. 311—355.)
- XX. Zur Geometrie der Zahlen 43
 (Verhandlungen des III. Internationalen Mathematiker-Kongresses,
 Heidelberg 1904, S. 164—173.)
- XXI. Diskontinuitätsbereich für arithmetische Äquivalenz 53
 (Crelles Journal für die reine und angewandte Mathematik, Bd. 129,
 S. 220—274; 1905.)

Zur Geometrie.

- XXII. Allgemeine Lehrsätze über die konvexen Polyeder. 103
 (Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen,
 mathematisch-physikalische Klasse, 1897, S. 198—219.)
- XXIII. Über die Begriffe Länge, Oberfläche und Volumen. 122
 (Jahresbericht der Deutschen Mathematiker-Vereinigung, Bd. 9, S. 115
 —121; 1901.)
- XXIV. Über die geschlossenen konvexen Flächen 128
 (In französischer Sprache unter dem Titel: „Sur les surfaces con-
 vexes fermées“ in den Comptes rendus de l'Académie des Sciences,
 Paris, t. 132, pp. 21—24; 1901.)
- XXV. Theorie der konvexen Körper, insbesondere Begründung ihres
 Oberflächenbegriffs 131
 (Bisher unveröffentlicht.)
- XXVI. Volumen und Oberfläche. 230
 (Mathematische Annalen, Bd. 57, S. 447—495; 1903.)
- XXVII. Über die Körper konstanter Breite 277
 (In russischer Sprache erschienen in: Mathematische Sammlung
 (Matematičeskij Sbornik), Moskau, Bd. 25, S. 505—508; 1904—1906.)

Zur Physik.

- XXVIII. Über die Bewegung eines festen Körpers in einer Flüssigkeit 283
 (Sitzungsberichte der K. Preußischen Akademie der Wissenschaften
 zu Berlin, Bd. XL, S. 1095—1110; 1888.)
- XXIX. Kapillarität. 298
 (Enzyklopädie der mathematischen Wissenschaften, V 1, Heft 4,
 S. 558—613.)

	Seite
XXX. Die Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern	352
(Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen, mathematisch-physikalische Klasse, 1908, S. 53—111.)	
XXXI. Eine Ableitung der Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern vom Standpunkte der Elektronentheorie	405
(Aus dem Nachlaß von Hermann Minkowski bearbeitet von Max Born. Mathematische Annalen, Bd. 68, S. 526—551; 1910.)	
XXXII. Raum und Zeit	431
(Physikalische Zeitschrift, 10. Jahrgang, S. 104—111, 1909 und Jahresbericht der Deutschen Mathematiker-Vereinigung, Bd. 18, S. 75—88, 1909; Sonderabdruck bei B. G. Teubner, Leipzig 1909.)	
Rede auf Dirichlet.	
XXXIII. Peter Gustav Lejeune Dirichlet und seine Bedeutung für die heutige Mathematik	447
(Jahresbericht der Deutschen Mathematiker-Vereinigung, Bd. 14, S. 149—163; 1905.)	
Sachregister	462
Berichtigungen	466

ZUR GEOMETRIE DER ZAHLEN

FORTSETZUNG

XIX.

Dichteste gitterförmige Lagerung kongruenter Körper.

(Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen.
Mathematisch-physikalische Klasse. 1904. S. 311—355.)

(Vorgelegt in der Sitzung vom 25. Juni 1904.)

Wir beschäftigen uns hier mit folgendem Probleme: *Gegeben ist ein beliebiger Grundkörper K . Lauter mit K kongruente und parallel orientierte Körper in unendlicher Anzahl sollen im Raume in gitterförmiger Anordnung derart gelagert werden, daß keine zwei der Körper ineinander eindringen und daß dabei der von den Körpern erfüllte Teil des Raumes gegenüber dem von ihnen freigelassenen möglichst groß ist.*

Unter einer *gitterförmigen Anordnung* der Körper verstehen wir, daß entsprechende Punkte in ihnen ein parallelepipedisches Punktsystem (Gitter) bilden. Wir beschränken die Untersuchung auf konvexe Körper.

Die Lösung dieses Problems gestattet interessante Anwendungen in der Zahlenlehre und ist auch für die Theorie von der Struktur der Kristalle von Bedeutung*). Aus der Kenntnis der dichtesten Lagerung von *Kugeln* folgen (§ 7) fast unmittelbar alle Tatsachen der Gauß-Dirichletschen Theorie der Parallelgitter (d. s. die Sätze über die arithmetische Reduktion der positiven ternären quadratischen Formen). Die dichteste Lagerung von *Oktaedern* (§ 9) gibt Aufschlüsse über die simultane Approximation zweier Größen durch rationale Zahlen mit gleichem Nenner. Die hier abgeleiteten allgemeinen Theoreme über gewisse Ungleichungen (§ 5) endlich bilden in ihrer Anwendung auf Parallelepipede und Oktaeder, bzw. auf Kreiszylinder und Doppelkegel die Grundlage für die zweckmäßigsten Algorithmen zur Ermittlung der Fundamenteinheiten in den kubischen Zahlkörpern von positiver bzw. negativer Diskriminante.

*) Vgl. in letzterer Hinsicht: Lord Kelvin, Baltimore lectures on molecular dynamics, London 1904, S. 618 ff.

§ 1. Analytische Formulierung des Problems und Reduktion auf den Fall von Körpern mit Mittelpunkt.

1. Es sei K ein beliebiger konvexer Körper. Wir führen rechtwinklige Koordinaten ξ, η, ζ mit einem Punkte O als Anfangspunkt ein, den wir uns im Inneren von K denken. Es sei J das Volumen von K . Sind λ, μ, ν irgendwelche Werte, so werde der größte Wert des linearen Ausdrucks $\lambda\xi + \mu\eta + \nu\zeta$ für die Punkte ξ, η, ζ im ganzen Bereiche von K mit $H(\lambda, \mu, \nu)$ bezeichnet. Diese Funktion H definiert den Körper K vollständig, sie heißt die *Stützebenenfunktion* von K . Sie genügt den Bedingungen

$$\left. \begin{aligned} H(0, 0, 0) &= 0, \quad H(\lambda, \mu, \nu) > 0 \\ H(t\lambda, t\mu, t\nu) &= tH(\lambda, \mu, \nu) \end{aligned} \right\} (\lambda, \mu, \nu \neq 0, 0, 0; t > 0),$$

$$H(\lambda_1 + \lambda_2, \mu_1 + \mu_2, \nu_1 + \nu_2) \leq H(\lambda_1, \mu_1, \nu_1) + H(\lambda_2, \mu_2, \nu_2),$$

und jede Funktion $H(\lambda, \mu, \nu)$, die diesen Bedingungen genügt, ist die Stützebenenfunktion eines konvexen Körpers mit O im Inneren*).

2. Bedeutet P einen Punkt ξ, η, ζ , so bezeichnen wir den Punkt $-\xi, -\eta, -\zeta$ als *Gegenpunkt* von P und mit P' .

Die Gegenpunkte zu den sämtlichen Punkten von K bilden den konvexen Körper K' mit der Stützebenenfunktion

$$H'(\lambda, \mu, \nu) = H(-\lambda, -\mu, -\nu),$$

das Spiegelbild K in bezug auf den Punkt O .

Die Funktion $\frac{1}{2}(H(\lambda, \mu, \nu) + H(-\lambda, -\mu, -\nu))$ bildet alsdann die Stützebenenfunktion für einen gewissen konvexen Körper, den wir mit $\frac{1}{2}(K + K')$ bezeichnen und der in O einen *Mittelpunkt* hat. Dieser Körper ist der Bereich aller solchen Punkte, welche irgendwie als Mitte einer Verbindungsstrecke eines Punktes von K mit einem Punkte von K' auftreten.

Ist z. B. K ein *Tetraeder*, so wird $\frac{1}{2}(K + K')$ ein *Oктаeder* mit Flächen parallel den Flächen des Tetraeders.

Das Volumen von $\frac{1}{2}(K + K')$ ist stets größer als das Volumen von K , wenn K ein Körper ohne Mittelpunkt ist (l. c. § 7). Hat K selbst einen Mittelpunkt, so entsteht K' und weiter $\frac{1}{2}(K + K')$ durch bloße Translation aus K .

3. Es sei

$$(1) \quad \xi = \alpha_1 x + \alpha_2 y + \alpha_3 z, \quad \eta = \beta_1 x + \beta_2 y + \beta_3 z, \quad \zeta = \gamma_1 x + \gamma_2 y + \gamma_3 z$$

*) Mathematische Annalen, Bd. 57, S. 447. Diese Ges. Abhandlungen, Bd. II, S. 230.

eine beliebige lineare Substitution mit einer von Null verschiedenen Determinante, deren Betrag Δ heie. Wir betrachten das *Zahlengitter* in x, y, z , d. i. das *parallelepipeditische* System $(\mathfrak{G})$ aller derjenigen Punkte G , für welche die Bestimmungsstücke x, y, z sämtlich *ganze Zahlen* sind. Den Bereich

$$0 \leq x < 1, \quad 0 \leq y < 1, \quad 0 \leq z < 1$$

nennen wir das *Grundparallelepipet* des Gitters $(\mathfrak{G})$, sein Volumen ist Δ . Durch Parallelverschiebung dieses Bereichs von O nach jedem einzelnen Punkte des Gitters erhalten wir eine lückenlose einfache Überdeckung des Raumes.

Wir denken uns ferner mit dem Körper K diese *Translationen* des Gitters, die Parallelverschiebungen vom Nullpunkte O nach den einzelnen Gitterpunkten G ausgeführt, und *wir nehmen an, daß alle so entstehenden Körper K_G (den Körper $K = K_0$ inbegriffen) gesondert sind, d. h. untereinander höchstens in den Begrenzungen zusammenstoßen.*

Zu diesem Ende wird bereits hinreichend sein, daß speziell der Körper K in keinen anderen der Körper K_G eindringt. Ist nun G ein beliebiger, von O verschiedener Gitterpunkt, so muß es dann irgendeine Ebene geben, welche K von K_G sondert so, daß die ihnen etwa gemeinsamen Punkte in die Ebene fallen und im übrigen K ganz auf einer Seite, K_G ganz auf der anderen Seite dieser Ebene bleibt. Sind λ, μ, ν die Richtungskosinus des von O auf die Ebene gefällten Lotes, so hat die Ebene von O einen Abstand $\geq H(\lambda, \mu, \nu)$ und von G einen Abstand $\geq H(-\lambda, -\mu, -\nu)$; die dazu parallele Ebene durch G hat daher von O einen Abstand

$$\geq H(\lambda, \mu, \nu) + H(-\lambda, -\mu, -\nu).$$

Der Gitterpunkt G liegt also außerhalb oder auf der Grenze desjenigen Körpers $\mathfrak{R} = K + K'$, der durch Dilatation des Körpers $\frac{1}{2}(K + K')$ vom Mittelpunkte O aus im Verhältnis 2:1 entsteht. Daraus entnehmen wir insbesondere die Folgerung:

Aus einem konvexen Körper K entstehen durch die Translationen eines Gitters $(\mathfrak{G})$ lauter gesonderte Körper dann und nur dann, wenn die entsprechende Tatsache für den zu K gehörigen konvexen Körper $\frac{1}{2}(K + K')$ mit Mittelpunkt statthat.

4. Die Begrenzung von $\mathfrak{R} = K + K'$ bezeichnen wir mit $\mathfrak{F}$. Wir definieren eine Funktion $\varphi(\xi, \eta, \zeta)$ für beliebige reelle Argumente, indem wir für die Punkte ξ, η, ζ auf der Fläche $\mathfrak{F}$ die Gleichung $\varphi(\xi, \eta, \zeta) = 1$, ferner allgemein

$$\varphi(t\xi, t\eta, t\zeta) = t\varphi(\xi, \eta, \zeta)$$

bei positivem t und $\varphi(0, 0, 0) = 0$ festsetzen.

Daß $\mathfrak{R}$ ein konvexer Körper ist, findet in der allgemeinen Funktionalungleichung

$$(2) \quad \varphi(\xi_1 + \xi_2, \eta_1 + \eta_2, \xi_1 + \xi_2) \leq \varphi(\xi_1, \eta_1, \xi_1) + \varphi(\xi_2, \eta_2, \xi_2),$$

und daß $\mathfrak{R}$ in O einen Mittelpunkt hat, in der Funktionalgleichung

$$\varphi(-\xi, -\eta, -\xi) = \varphi(\xi, \eta, \xi)$$

seinen Ausdruck.

Wir setzen ferner mit Rücksicht auf die Substitution (1):

$$\varphi(\xi, \eta, \xi) = f(x, y, z).$$

Damit alle Körper K_G gesondert liegen, muß alsdann für jedes von $0, 0, 0$ verschiedene ganzzahlige System x, y, z die Ungleichung

$$(3) \quad f(x, y, z) \geq 1 \quad (x, y, z \neq 0, 0, 0)$$

bestehen.

5. Im unendlichen Raume ist nunmehr einem jeden Gitterpunkt G einerseits genau ein Volumen $= \Delta$ und andererseits ein mit K kongruenter Körper vom Volumen J zugeordnet, also muß jedenfalls

$$(4) \quad J \leq \Delta$$

sein, und wir werden sagen können, der von den Körpern K_G erfüllte Raum verhält sich zu dem ganzen unendlichen Raume wie $J : \Delta$.

Wir bemerken insbesondere, daß die Körper K_G den Raum lückenlos erfüllen, wenn $J = \Delta$ ist. Da nun das Volumen J^* von $\frac{1}{2}(K + K')$ ebenfalls noch $\leq \Delta$, aber, wenn K keinen Mittelpunkt hat, $> J$ ist, so schließen wir, daß eine lückenlose Überdeckung des Raumes durch die Körper K_G notwendig das Vorhandensein eines Mittelpunktes in K verlangt.

Die Aufgabe, die wir uns zu stellen haben, ist nun, die Koeffizienten $\alpha_1, \dots, \gamma_3$ der Substitution (1) derart zu wählen, daß, während die sämtlichen Ungleichungen (3) erfüllt sind, dabei der Betrag der Determinante Δ möglichst klein ausfalle.

§ 2. Hilfssatz über ganzzahlige Substitutionen.

6. Es seien $\mathfrak{A}, \mathfrak{B}, \mathfrak{C}$ irgend drei von O verschiedene Punkte des Gitters in x, y, z , welche in drei *unabhängigen*, d. h. nicht in eine Ebene fallenden Richtungen von O aus liegen, so entsteht die Frage, wie man von den vier Punkten $O, \mathfrak{A}, \mathfrak{B}, \mathfrak{C}$ aus dieses ganze Gitter ($\mathfrak{G}$) aufbauen kann.

Zu jedem Punkte $P(x, y, z)$ im Raume gehören bestimmte Werte $\mathfrak{x}, \mathfrak{y}, \mathfrak{z}$, so daß der Vektor OP mit den Vektoren $O\mathfrak{A}, O\mathfrak{B}, O\mathfrak{C}$ durch eine Relation

$$OP = \mathfrak{x} \cdot O\mathfrak{A} + \mathfrak{y} \cdot O\mathfrak{B} + \mathfrak{z} \cdot O\mathfrak{C}$$

verbunden ist, und sind danach x, y, z gewisse lineare Formen in ξ, η, ζ mit *ganzzahligen* Koeffizienten. Den Bereich

$$0 \leq \xi < 1, \quad 0 \leq \eta < 1, \quad 0 \leq \zeta < 1$$

bezeichnen wir kurz als das Parallelepiped $O(\mathfrak{A}\mathfrak{B}\mathfrak{C})$ und die Bestimmungsstücke ξ, η, ζ nennen wir die *zu diesem Parallelepiped gehörigen* Koordinaten.

Das Gitter ($\mathfrak{G}$) besitzt diese *Grundeigenschaft*: Sind G_0, G_1, G_2 irgend drei Punkte des Gitters, so gehört derjenige Punkt G_3 , für den Vektor $G_2 G_3 = G_0 G_1$ ist, stets ebenfalls zum Gitter. Auf Grund dieses *parallelogrammatischen* Charakters des Gitters können wir durch die folgenden Forderungen drei gewisse Gitterpunkte A, B, C völlig eindeutig fixieren (Fig. 1):

1) es soll A von O verschieden sein und auf der Strecke $O\mathfrak{A}$ möglichst nahe an O liegen;

2) es soll B in der Ebene $O\mathfrak{A}\mathfrak{B}$ außerhalb der Geraden $O\mathfrak{A}$ auf derselben Seite von ihr wie $\mathfrak{B}$ zunächst möglichst nahe an dieser Geraden und nächstdem derart liegen, daß von O nach einem Punkte des Strahles $O\mathfrak{B}$ ein Vektor

$$OB + X \cdot OA$$

hinführt, wobei $0 \leq X < 1$ ist;

3) es soll C außerhalb der Ebene $O\mathfrak{A}\mathfrak{B}$ nach derselben Seite hin wie $\mathfrak{C}$ zunächst möglichst nahe an dieser Ebene und nächstdem so liegen, daß von O nach einem Punkte des Strahles $O\mathfrak{C}$ ein Vektor

$$OC + X \cdot OA + Y \cdot OB$$

hinführt, wobei $0 \leq X < 1, 0 \leq Y < 1$ ist.

7. Nachdem A, B, C ermittelt sind, haben wir für jeden Punkt P im Raume eine bestimmte Vektorenrelation:

$$OP = X \cdot OA + Y \cdot OB + Z \cdot OC.$$

Bei ganzzahligen X, Y, Z finden wir in P jedesmal einen Gitterpunkt aus ($\mathfrak{G}$) und sind infolgedessen x, y, z gleich linearen *ganzzahligen* Formen in X, Y, Z . Nach der Bestimmungsweise von A, B, C aber kann es keinen Gitterpunkt x, y, z außer dem Nullpunkte geben, für den $0 \leq Z < 1, 0 \leq Y < 1, 0 \leq X < 1$ ist, und gibt es infolgedessen überhaupt keinen Gitterpunkt x, y, z außer denjenigen, für welche auch X, Y, Z ganzzahlige Werte haben, d. h. auch X, Y, Z sind lineare *ganzzahlige* Formen in x, y, z . Mithin ist nicht bloß die Determinante der Formen x, y, z in X, Y, Z , sondern auch deren reziproker Wert eine ganze Zahl

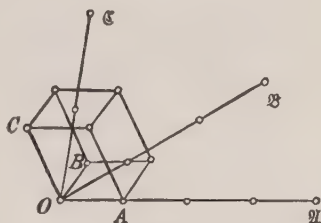


Fig. 1.

und also diese Determinante ± 1 , mithin das Gitter ($\mathfrak{G}$) völlig identisch mit dem Zahlengitter in X, Y, Z . Wir kommen damit zu folgendem Satze:

Sind x, y, z irgendwelche lineare ganzzahlige Formen in $\mathfrak{x}, \mathfrak{y}, \mathfrak{z}$ mit einer von Null verschiedenen Determinante, so kann man stets (und nur auf eine Weise) drei lineare ganzzahlige Formen X, Y, Z in x, y, z mit einer Determinante ± 1 einführen, daß Relationen

$$X = a_1 \mathfrak{x} + a_2 \mathfrak{y} + a_3 \mathfrak{z}, \quad Y = b_1 \mathfrak{x} + b_2 \mathfrak{y} + b_3 \mathfrak{z}, \quad Z = c_1 \mathfrak{x} + c_2 \mathfrak{y} + c_3 \mathfrak{z}$$

entstehen, wobei die ganzzahligen Koeffizienten den Bedingungen

$$a_1 > 0, \quad b_1 > 0, \quad c_1 > 0; \quad 0 \leq \frac{a_2}{b_1}, \quad \frac{a_3}{c_1}, \quad \frac{b_2}{c_1} < 1$$

genügen.

Die Einführung der Variablen X, Y, Z anstatt x, y, z nennen wir die *Reduktion* des Zahlengitters ($\mathfrak{G}$) in bezug auf das Parallelepiped $O(\mathfrak{A}\mathfrak{B}\mathfrak{C})$ und ferner bezeichnen wir $O(ABC)$ als *reduziertes Grundparallelepiped*.

§ 3. Zusammenrücken einer gitterförmigen Lagerung in drei Richtungen.

8. Wir denken uns die 9 Koeffizienten $\alpha_1, \dots, \gamma_3$ in (1), welche ein Gitter ($\mathfrak{G}$) in bezug auf den Körper K orientieren, als stetig veränderliche Größen, während alle Ungleichungen (3) bestehen sollen, also kein Gitterpunkt außer O ins Innere des Bereichs $\mathfrak{R} = K + K'$ falle. Nach (4) muß dabei stets $\Delta \geq J$ bleiben. Wir wollen nun zunächst zeigen:

Ein Minimum von Δ kann jedenfalls nur eintreten, während irgend drei Gitterpunkte in drei unabhängigen Richtungen von O aus auf die Begrenzung von $\mathfrak{R}$ zu liegen kommen.

In der Tat, nehmen wir an, auf der Begrenzung $\mathfrak{F}$ von $\mathfrak{R}$ liegen entweder a) gar keine Gitterpunkte oder b) nur ein Paar Gitterpunkte an den Enden eines Durchmessers durch O oder c) mehrere Paare von Gitterpunkten, aber nur in einer einzigen Ebene durch O . Wir wählen dann drei Gitterpunkte $\mathfrak{A}, \mathfrak{B}, \mathfrak{C}$ aus, die nicht mit O in einer Ebene liegen, entweder im Falle a) beliebig oder im Falle b) so, daß $\mathfrak{A}$ auf $\mathfrak{F}$ liegt, oder im Falle c) so, daß $\mathfrak{A}$ und $\mathfrak{B}$ auf $\mathfrak{F}$ liegen. Wir führen nunmehr nach der in § 2 dargelegten Methode die Reduktion des Gitters ($\mathfrak{G}$) in bezug auf das Parallelepiped $O(\mathfrak{A}\mathfrak{B}\mathfrak{C})$ aus, und es sei $O(ABC)$ das neu einzuführende reduzierte Grundparallelepiped.

Wir variieren die Substitution (1) in der Weise, daß wir A und B festhalten, dagegen C geradlinig auf den Punkt O zurücklassen. Dabei verringert sich der Wert von Δ , des Volumens des Grundparallelepiped, während die auf $\mathfrak{F}$ bereits vorhandenen Gitterpunkte völlig unberührt bleiben. Mit Rücksicht auf die Bedingung $\Delta \geq J$ wird die Variation

schließlich auf einen Zustand auslaufen müssen, wobei irgendein weiterer Gitterpunkt außerhalb der Ebene OAB auf $\mathfrak{F}$ auftritt.

Da dieses Verfahren sich bei den dreierlei Umständen a), b), c) ausführen läßt, so können wir in der Tat das Gitter ($\mathfrak{G}$) unter Erhaltung der Relationen (3) und kontinuierlicher Verringerung von Δ so lange variieren, bis wir auf $\mathfrak{F}$ drei Gitterpunkte $\mathfrak{A}$, $\mathfrak{B}$, $\mathfrak{C}$ in drei unabhängigen Richtungen von O aus vorfinden.

9. Mit $\mathfrak{A}$, $\mathfrak{B}$, $\mathfrak{C}$ zugleich liegen auch deren Gegenpunkte $\mathfrak{A}'$, $\mathfrak{B}'$, $\mathfrak{C}'$ auf $\mathfrak{F}$, und mit diesen sechs Punkten enthält $\mathfrak{K}$ als konvexer Körper den ganzen Bereich des Oktaeders $\mathfrak{ABC}\mathfrak{A}'\mathfrak{B}'\mathfrak{C}'$, welches diese sechs Punkte als Ecken hat. Wir wollen dieses Oktaeder kurz *das durch $\mathfrak{A}$, $\mathfrak{B}$, $\mathfrak{C}$ bestimmte Oktaeder* nennen und mit $\text{Okt}(\mathfrak{ABC})$ bezeichnen.

Enthält dieses Oktaeder außer O und den Ecken noch irgendeinen weiteren Gitterpunkt G , so gehört dieser auch zu $\mathfrak{K}$ und befindet sich daher notwendig ebenfalls auf der Fläche $\mathfrak{F}$. Es liege G etwa außerhalb der Ebene $O\mathfrak{AB}$, so wird das $\text{Okt}(\mathfrak{ABG})$ von kleinerem Volumen als das $\text{Okt}(\mathfrak{ABC})$ sein. Daraus entnehmen wir das Resultat:

Werden unter allen auf der Begrenzung von $\mathfrak{K}$ irgend vorhandenen Gitterpunkten drei nicht mit O in einer Ebene gelegene Punkte $\mathfrak{A}$, $\mathfrak{B}$, $\mathfrak{C}$ derart ausgewählt, daß das $\text{Okt}(\mathfrak{ABC})$ ein möglichst kleines Volumen hat, so enthält dieses Oktaeder gewiß keinen Gitterpunkt außer den Ecken und dem Mittelpunkt.

§ 4. Gitteroktaeder.

10. Ein Oktaeder $\mathfrak{ABC}\mathfrak{A}'\mathfrak{B}'\mathfrak{C}'$ mit O als Mittelpunkt, dessen sechs Ecken Gitterpunkte sind, (die nicht in eine Ebene fallen), und welches außer O und den Ecken keinen Gitterpunkt enthält, soll ein *Gitteroktaeder* heißen. Wir fragen jetzt, wie können wir von einem solchen Oktaeder aus das ganze gegebene Gitter ($\mathfrak{G}$) herstellen.

Zu jedem Punkte $P(x, y, z)$ gehören bestimmte Größen $\mathfrak{x}$, $\mathfrak{y}$, $\mathfrak{z}$, mit denen

$$OP = \mathfrak{x} \cdot O\mathfrak{A} + \mathfrak{y} \cdot O\mathfrak{B} + \mathfrak{z} \cdot O\mathfrak{C}$$

gilt. Der Bereich des $\text{Okt}(\mathfrak{ABC})$ ist dann durch

$$(5) \quad |\mathfrak{x}| + |\mathfrak{y}| + |\mathfrak{z}| \leq 1$$

definiert.

Wir führen nach § 2 die Reduktion des Zahlengitters in bezug auf das Parallelepiped $O(\mathfrak{ABC})$ aus. Es sei $O(ABC)$ das reduzierte Grundparallelepiped und X, Y, Z seien die dazu gehörigen neuen Gitterkoordinaten, für welche sich Relationen

$$(6) \quad X = a_1\mathfrak{x} + a_2\mathfrak{y} + a_3\mathfrak{z}, \quad Y = b_1\mathfrak{x} + b_2\mathfrak{y} + b_3\mathfrak{z}, \quad Z = c_1\mathfrak{x} + c_2\mathfrak{y} + c_3\mathfrak{z}$$

mit den Bedingungen

$$a_1 > 0; \quad b_2 > 0, \quad 0 \leq \frac{a_2}{b_2} < 1; \quad c_3 > 0, \quad 0 \leq \frac{a_3}{c_3}, \quad \frac{b_3}{c_3} < 1$$

herausstellen, während das vorgelegte Gitter ($\mathfrak{G}$) sich genau mit den sämtlichen ganzzahligen Systemen X, Y, Z deckt. Für $\mathfrak{A}, \mathfrak{B}, \mathfrak{C}$ haben wir bzw.

$$X, Y, Z = a_1, 0, 0; \quad a_2, b_2, 0; \quad a_3, b_3, c_3.$$

Da innerhalb der Strecke $O\mathfrak{A}$ im Gitteroktaeder kein Gitterpunkt liegt, so finden wir zunächst $A = \mathfrak{A}$, $a_1 = 1$.

Sodann lautet die Ungleichung (5) für $z = 0$, also in der Ebene $O\mathfrak{A}\mathfrak{B}$:

$$\left| X - \frac{a_2}{b_2} Y \right| + \left| \frac{Y}{b_2} \right| \leq 1.$$

Wenn nun $b_2 > 1$, also ≥ 2 wäre, so könnten wir $Y = 1$, und, da $0 \leq a_2 < b_2$ ist, die Größe $X = 0$ bzw. $= 1$ derart annehmen, daß $\left| X - \frac{a_2}{b_2} \right| \leq \frac{1}{2}$ ist; dann würde die Ungleichung hier erfüllt sein, und hätten wir damit in der Ebene $O\mathfrak{A}\mathfrak{B}$ einen von $O, \mathfrak{A}, \mathfrak{A}', \mathfrak{B}, \mathfrak{B}'$ verschiedenen und dem Oktaeder angehörigen Gitterpunkt gefunden, was der Natur eines Gitteroktaeders widerspricht. Also muß $b_2 = 1$, $a_2 = 0$ sein, und wir finden demnach $B = \mathfrak{B}$.

Die Ungleichung (5) lautet nunmehr

$$(7) \quad \left| X - \frac{a_3}{c_3} Z \right| + \left| Y - \frac{b_3}{c_3} Z \right| + \left| \frac{Z}{c_3} \right| \leq 1.$$

Wir unterscheiden mehrere Fälle. Ist $c_3 > 1$ und *ungerade* $= 2d + 1$, so setzen wir $Z = 1$ und können $X = 0$ bzw. 1 und $Y = 0$ bzw. 1 so annehmen, daß $\left| X - \frac{a_3}{c_3} \right| < \frac{1}{2}$, $\left| Y - \frac{b_3}{c_3} \right| < \frac{1}{2}$ ist. Dann wird die Summe links in der Ungleichung (7)

$$\leq \frac{d}{2d+1} + \frac{d}{2d+1} + \frac{1}{2d+1} = 1;$$

wir hätten also einen von O und den Ecken verschiedenen Gitterpunkt im Okt($\mathfrak{A}\mathfrak{B}\mathfrak{C}$) gefunden, was nach Voraussetzung nicht sein darf.

Ist $c_3 > 1$ und *gerade* $= 2d$, so können wir analog $Z = 1$ und $X = 0$ bzw. 1 , $Y = 0$ bzw. 1 derart wählen, daß $\left| X - \frac{a_3}{c_3} \right| \leq \frac{1}{2}$, $\left| Y - \frac{b_3}{c_3} \right| \leq \frac{1}{2}$ ist. Dabei wird, wofern nicht in diesen beiden Ungleichungen zugleich das Gleichheitszeichen gilt, die Summe links in (7):

$$\leq \frac{(d-1) + d + 1}{2d} = 1$$

ausfallen, und wäre Okt($\mathfrak{A}\mathfrak{B}\mathfrak{C}$) wieder nicht ein Gitteroktaeder. In dem eben erwähnten besonderen Falle aber hätten wir $a_3 = d$, $b_3 = d$, $c_3 = 2d$

und, wenn $d > 1$ wäre, würde der Gitterpunkt $X, Y, Z = 1, 1, 2$ innerhalb der Strecke $O\mathfrak{C}$, also im Inneren des Okt($\mathfrak{A}\mathfrak{B}\mathfrak{C}$) liegen.

Danach erkennen wir, daß für a_3, b_3, c_3 allein die beiden Wertsysteme $0, 0, 1$ oder $1, 1, 2$ in Frage kommen.

In dem ersteren Falle werden die Gleichungen (6)

$$X = \mathfrak{x}, \quad Y = \mathfrak{y}, \quad Z = \mathfrak{z}$$

und ist das ursprüngliche Zahlengitter ($\mathfrak{G}$) in x, y, z völlig identisch mit dem Zahlengitter in $\mathfrak{x}, \mathfrak{y}, \mathfrak{z}$. In diesem Falle bezeichnen wir das Okt($\mathfrak{A}\mathfrak{B}\mathfrak{C}$) als *Gitteroktaeder erster Art*.

Im zweiten Falle haben wir

$$X = \mathfrak{x} + \mathfrak{z}, \quad Y = \mathfrak{y} + \mathfrak{z}, \quad Z = 2\mathfrak{z}.$$

Ganzzahlige Werte X, Y, Z entstehen einmal, wenn $\mathfrak{x}, \mathfrak{y}, \mathfrak{z}$ ganze Zahlen sind, und ferner, wenn $\mathfrak{x} - \frac{1}{2}, \mathfrak{y} - \frac{1}{2}, \mathfrak{z} - \frac{1}{2}$ ganze Zahlen sind. Insbesondere ist der Punkt C^* :

$$\mathfrak{x} = \frac{1}{2}, \quad \mathfrak{y} = \frac{1}{2}, \quad \mathfrak{z} = \frac{1}{2},$$

d. i. der Mittelpunkt des Parallelepipeds $O(\mathfrak{A}\mathfrak{B}\mathfrak{C})$, ein Punkt des Gitters in x, y, z . Das ganze ursprüngliche Gitter ($\mathfrak{G}$) besteht hier einmal aus dem Gitter in $\mathfrak{x}, \mathfrak{y}, \mathfrak{z}$ und sodann aus den Punkten, die wir aus dem letzteren Gitter durch die Translation von O nach C^* ableiten. In diesem Falle bezeichnen wir das Okt($\mathfrak{A}\mathfrak{B}\mathfrak{C}$) als *Gitteroktaeder zweiter Art*. Das Volumen des Parallelepipeds $O(ABC)$ ist hier die Hälfte des Volumens von $O(\mathfrak{A}\mathfrak{B}\mathfrak{C})$.

11. Wir können nunmehr folgenden Satz zur Charakterisierung der Gitteroktaeder aussprechen:

Sind $x_1, y_1, z_1; x_2, y_2, z_2; x_3, y_3, z_3$ die Koordinaten von drei Gitterpunkten $\mathfrak{A}, \mathfrak{B}, \mathfrak{C}$, so ist das Oktaeder mit den Ecken $\mathfrak{A}, \mathfrak{B}, \mathfrak{C}, \mathfrak{A}', \mathfrak{B}', \mathfrak{C}'$ ein Gitteroktaeder erster Art, wenn die Determinante aus jenen Koordinaten ± 1 ist, und ein Gitteroktaeder zweiter Art, wenn jene Determinante ± 2 ist und zudem

$$\frac{1}{2}(x_1 + x_2 + x_3), \quad \frac{1}{2}(y_1 + y_2 + y_3), \quad \frac{1}{2}(z_1 + z_2 + z_3)$$

gleich ganzen Zahlen sind.

12. Wir fügen noch folgende Bemerkung hinzu. Die linke Seite in (7) wird, wenn wir uns $0 \leq a_3 \leq b_3 \leq c_3$ und $c_3 \geq 2$ denken, sogar < 1 durch wenigstens eines der Systeme $X, Y, Z = 0, 0, 1; 0, 1, 1; 1, 1, 1$, es sei denn, daß $c_3 = 2d + 1$ und dazu $a_3, b_3 = d, d; d, d + 1; d + 1, d + 1$ oder $c_3 = 2d$ und dazu $a_3, b_3 = d - 1, d; d, d + 1$ bzw. $= d, d$ ist. In den letzteren Fällen hat die linke Seite von (7) für das System

$X, Y, Z = 1, 1, 2$ den Wert $\frac{4}{2d+1}$ (< 1 für $c_3 > 3$) oder $\frac{4}{2d}$ (< 1 für $c_3 > 4$) bzw. $\frac{2}{2d}$ (< 1 für $c_3 > 2$). Wir finden danach unter den eben genannten vier Systemen X, Y, Z immer wenigstens ein solches, wofür die linke Seite von (7) sich < 1 erweist, außer wenn

$a_3, b_3, c_3 = 0, 1, 2; 1, 1, 2; 1, 2, 2; 1, 1, 3; 1, 2, 3; 2, 2, 3; 1, 2, 4; 2, 3, 4$ ist.

§ 5. Reduktion der unendlich vielen Ungleichungen des Problems auf eine endliche Anzahl.

13. Wir nehmen wieder an, daß kein Gitterpunkt außer O im Inneren von $\mathfrak{R}$ liegt, daß aber auf der Begrenzung $\mathfrak{F}$ von $\mathfrak{R}$ drei Gitterpunkte A, B, C — wir ändern damit etwas die Bezeichnung — vorhanden sind, die ein Gitteroktaeder erster oder zweiter Art bestimmen. Wir führen die Koordinaten X, Y, Z zum Parallelepiped $O(ABC)$ ein und setzen die in 4. für $\mathfrak{R}$ definierte Funktion

$$\varphi(\xi, \eta, \xi) = F(X, Y, Z).$$

Die dort angegebenen Funktionalbedingungen gehen dann in

$$F(tX, tY, tZ) = tF(X, Y, Z), \quad t > 0,$$

$$(8) \quad F(X_1, Y_1, Z_1) + F(X_2, Y_2, Z_2) \geq F(X_1 + X_2, Y_1 + Y_2, Z_1 + Z_2),$$

$$(9) \quad F(-X, -Y, -Z) = F(X, Y, Z)$$

über. Indem die Punkte A, B, C auf $\mathfrak{F}$ liegen, haben wir

$$(10) \quad F(1, 0, 0) = 1, \quad F(0, 1, 0) = 1, \quad F(0, 0, 1) = 1.$$

Weiter soll nun auf Grund unserer Voraussetzung die Ungleichung

$$(I) \quad F(X, Y, Z) \geq 1 \quad (X, Y, Z \neq 0, 0, 0)$$

für jedes von $0, 0, 0$ verschiedene ganzzahlige Wertsystem X, Y, Z gelten, und soll überdies, falls $\text{Okt}(ABC)$ ein Gitteroktaeder zweiter Art vorstellt, die Ungleichung

$$(II) \quad F\left(X + \frac{1}{2}, Y + \frac{1}{2}, Z + \frac{1}{2}\right) \geq 1$$

für jedes beliebige ganzzahlige System X, Y, Z (das System $0, 0, 0$ hier inbegriffen), gelten.

14. Wir wollen jetzt zeigen, daß infolge des Bestehens der drei Gleichungen (10) von allen diesen Ungleichungen gewisse in *endlicher Anzahl* vorhandene das Bestehen aller übrigen nach sich ziehen.

In der Tat, greifen wir von den Ungleichungen (I) zunächst diejenigen besonderen heraus, in denen die Zahlen X, Y, Z nur Werte 0 oder ± 1 haben, das sind die folgenden:

$$(11) \quad F(\pm 1, \pm 1, 0) \geq 1, \quad F(\pm 1, 0, \pm 1) \geq 1, \quad F(0, \pm 1, \pm 1) \geq 1, \\ (12) \quad F(\pm 1, \pm 1, \pm 1) \geq 1$$

für alle möglichen Werte der einzelnen Vorzeichen ± 1 .

Sollte nun, während alle Bedingungen (10), (11), (12) bereits statthaben, irgendeine der Ungleichungen (I) nicht gelten, so sei $G(X, Y, Z = a, b, c)$ ein Gitterpunkt, für den

$$F(a, b, c) < 1$$

ist; wir können uns etwa $0 \leq a \leq b \leq c$ ($c \geq 2$) vorstellen, da wir die Rollen von $\pm X, \pm Y, \pm Z$ vertauschen können. Das Okt(ABG) gehört dann mit allen Punkten, die außerhalb der Ebene OAB liegen, ganz zum Inneren von $\mathfrak{K}$. Dieses Oktaeder ist durch

$$\left| X - \frac{a}{c} Z \right| + \left| Y - \frac{b}{c} Z \right| + \left| \frac{Z}{c} \right| \leq 1$$

bestimmt. Ähnlich wie in 10. schließen wir nun, daß in diesem Oktaeder, falls nicht gerade a, b, c von der Form $d, d, 2d$ ist, stets wenigstens einer von den Gitterpunkten $Z = 1, X = 0$ oder $1, Y = 0$ oder 1 auftritt. Der betreffende Gitterpunkt dürfte aber nach einer der Ungleichungen (11), (12) nicht ins Innere von $\mathfrak{K}$ fallen. Mithin erkennen wir, daß zu den Ungleichungen (11) und (12) nur noch die folgenden

$$(13) \quad F(\pm 1, \pm 1, \pm 2) \geq 1, \quad F(\pm 1, \pm 2, \pm 1) \geq 1, \quad F(\pm 2, \pm 1, \pm 1) \geq 1$$

hinzuzunehmen sind, um die Gesamtheit der Ungleichungen (I) sicherzustellen. Wir haben damit den Satz gewonnen:

Ist Okt(ABC) ein Gitteroktaeder erster Art, so haben die Bedingungen (10), (11), (12), (13) die Gesamtheit der unendlich vielen Ungleichungen (I) zur Folge.

Ist das Okt(ABC) ein Gitteroktaeder zweiter Art, so haben wir unter den Ungleichungen (II), indem wir X, Y, Z die Werte 0 und -1 beilegen, insbesondere die Ungleichungen:

$$(14) \quad F\left(\pm \frac{1}{2}, \pm \frac{1}{2}, \pm \frac{1}{2}\right) \geq 1.$$

Wir werden nun zeigen:

Ist Okt(ABC) ein Gitteroktaeder zweiter Art, so reichen die Gleichungen (10) und die Ungleichungen (14) bereits aus, um die Gesamtheit der unendlich vielen Ungleichungen (I) und (II) nach sich zu ziehen.

Zunächst folgen aus (14) alle Ungleichungen (I), nämlich speziell die Ungleichungen (11), (12), (13). Denn wir schließen aus (14) mit Rücksicht auf (8):

$$F(1, 1, 1) \geq 2, \\ F(1, 1, 0) + F(0, 0, 1) \geq F(1, 1, 1) \geq 2, \quad F(1, 1, 0) \geq 1, \\ F(1, 1, 2) + F(0, 0, -1) \geq F(1, 1, 1), \quad F(1, 1, 2) \geq 1,$$

und ähnlich bei Änderung der Vorzeichen der einzelnen Variablen und bei den Permutationen ihrer Reihenfolge.

Was sodann die Ungleichungen (II) anbelangt, so nehmen wir an, es seien die Ungleichungen (14) erfüllt, es bestünde aber noch für einen Gitterpunkt $G\left(X, Y, Z = a + \frac{1}{2}, b + \frac{1}{2}, c + \frac{1}{2}\right)$, wobei a, b, c ganze Zahlen sind, die Ungleichung

$$F\left(\frac{2a+1}{2}, \frac{2b+1}{2}, \frac{2c+1}{2}\right) < 1;$$

wir können ohne wesentliche Beschränkung uns etwa

$$0 \leq \frac{2a+1}{2} \leq \frac{2b+1}{2} \leq \frac{2c+1}{2}$$

und $2c+1 \geq 3$ denken. Nun enthält das Okt (ABG) , d. i. der Bereich

$$\left|X - \frac{2a+1}{2c+1}Z\right| + \left|Y - \frac{2b+1}{2c+1}Z\right| + \left|\frac{2Z}{2c+1}\right| \leq 1$$

den Punkt $X = \frac{1}{2}, Y = \frac{1}{2}, Z = \frac{1}{2}$; denn für diesen Punkt wird die hier links stehende Summe

$$\frac{(c-a) + (c-b) + 1}{2c+1} \leq 1;$$

also müßte auch $F\left(\frac{1}{2}, \frac{1}{2}, \frac{1}{2}\right) < 1$ sein im Widerspruch mit einer der Ungleichungen (14). Die Ungleichungen (14) haben demnach in der Tat mit Rücksicht auf die drei Gleichungen (10) alle Ungleichungen (I) und (II) zur Folge.

15. Nach der Bemerkung in 12. finden wir, wenn für ein ganzzahliges System $G(a, b, c)$ die Umstände $0 \leq a \leq b \leq c$ und $c \geq 2$ statthaben und a, b, c nicht gerade mit einem der Systeme

$$0, 1, 2; 1, 1, 2; 1, 2, 2; 1, 1, 3; 1, 2, 3; 2, 2, 3; 1, 2, 4; 2, 3, 4$$

übereinstimmt, wenigstens einen von den Gitterpunkten $0, 0, 1; 0, 1, 1; 1, 1, 1; 1, 1, 2$ sogar im Inneren des Okt (ABG) gelegen und kann daher G auch nicht auf der Begrenzung von $\mathfrak{K}$ liegen, sondern muß dafür notwendig

$$F(a, b, c) > 1$$

ausfallen.

§ 6. Die benachbarten Körper in einer dichtesten Lagerung.

16. Nach den letzten Ausführungen ist klar, daß bei kontinuierlicher Variation der Koeffizienten $\alpha_1, \alpha_2, \dots, \gamma_3$ in (1) d. h. des Gitters ($\mathfrak{G}$) keine einzige der Ungleichungen (I) bzw. der Ungleichungen (I) und (II) verletzt wird, so lange nur die drei Gleichungen (10) und die in endlicher Anzahl vorhandenen Ungleichungen (11), (12), (13) bzw. (14) in Kraft

bleiben. Unter beständiger Wahrung dieser Bedingungen (10)—(13) bzw. (10), (14) suchen wir nun weiter $\alpha_1, \alpha_2, \dots, \gamma_3$ kontinuierlich derart zu variieren, daß Δ sich verringert oder wenigstens nicht zunimmt und dabei in möglichst vielen unabhängigen von den fraglichen Ungleichungen das Gleichheitszeichen sich einstellt.

17. Wir nehmen als ersten Fall an, daß Okt (ABC) ein Gitteroktaeder erster Art ist und können dazu voraussetzen, daß auch im Verlaufe der vorzunehmenden Variationen niemals drei Gitterpunkte auf $\mathfrak{F}$ auftreten sollen, die ein Gitteroktaeder zweiter Art bestimmen.

Zunächst möge in keiner der Ungleichungen (11), (12), (13) das Gleichheitszeichen gelten. Wir projizieren den ganzen Bereich des Körpers $\mathfrak{K}$ in Richtung OC auf irgendeine diese Richtung nicht enthaltende Ebene durch O . Dadurch entsteht in dieser Ebene eine gewisse konvexe Figur $\mathfrak{P}$ mit O als Mittelpunkt; $\mathfrak{A}$ und $\mathfrak{B}$ seien darin die Projektionen von A und B (Fig. 2). Im allgemeinen wird es nun möglich sein, unter Festhaltung von B und C den Punkt A auf der Begrenzung von $\mathfrak{K}$ kontinuierlich so zu bewegen, daß währenddessen $\mathfrak{A}$ geradlinig auf O zuwandert und damit Δ kleiner wird. Eine solche Variation ist nur in dem Falle zunächst nicht ausführbar, wenn A innerhalb einer auf der Fläche $\mathfrak{F}$ verlaufenden, mit OC parallelen Strecke liegt, die sich dann natürlich in einen Randpunkt von $\mathfrak{P}$ projiziert. In diesem Falle würden wir A zuvörderst nach einem Endpunkte jener Strecke hin wandern lassen, wobei Δ sich nicht ändert, um hernach, wenn inzwischen kein neuer Gitterpunkt auf $\mathfrak{F}$ aufgetreten ist, Δ in der beschriebenen Weise zu verringern. Da nun Δ nicht unter die von vornherein gegebene Grenze J , das Volumen von K , sinken kann, muß der dargelegte Prozeß notwendig einmal auf einen Zustand auslaufen, wobei in einer der Ungleichungen (11), (12), (13) das Zeichen $=$ eintritt.

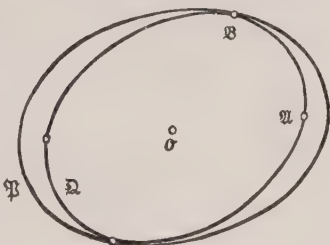


Fig. 2.

Wenn dabei z. B. $F(1, 1, 2) = 1$ würde, so fänden wir nach dem Satze aus 11. in $-1, 0, 0$; $0, -1, 0$; $1, 1, 2$ drei Gitterpunkte auf $\mathfrak{F}$, die ein Gitteroktaeder zweiter Art bestimmen, was wir ausgeschlossen haben. Hiernach kommt gegenwärtig keine der Ungleichungen (13) in Frage. Wir machen nun ferner zunächst die Annahme, daß bei den vorzunehmenden Variationen auch stets die Ungleichungen

$$F(\pm 1, \pm 1, \pm 1) > 1$$

gewahrt bleiben. Dann kommen einstweilen nur die Ungleichungen (11) für das Eintreten des Gleichheitszeichens in Betracht. Da wir die Be-

zeichnungen Y und Z , auch X mit $-X$ vertauschen können, so dürfen wir annehmen, daß der Gitterpunkt $S(-1, 0, 1)$ auf der Fläche $\mathfrak{F}$ erscheint, daß also

$$F(-1, 0, 1) = 1$$

wird. Diese Gleichung gewährleistet vermöge

$$F(1, 0, 1) + F(-1, 0, 1) \geq F(0, 0, 2) = 2$$

die Ungleichung $F(1, 0, 1) \geq 1$, d. h. solange C und S auf $\mathfrak{F}$ bleiben, kann $1, 0, 1$ gewiß nicht ins Innere von $\mathfrak{R}$ fallen.

Die weiteren von den Ungleichungen (11) mögen noch das Zeichen $>$ behalten haben. Wir können nun mit dem Punkte B operieren wie vorhin mit A , indem wir wieder die Parallelprojektion von $\mathfrak{R}$ in Richtung OC verwenden und können ohne Zunahme von Δ erzielen, daß in einer neuen der Ungleichungen (11) das Zeichen $=$ eintritt. Indem wir noch X und Z vertauschen, auch Y durch $-Y$ ersetzen können, dürfen wir annehmen, es trete jetzt der Gitterpunkt $R(0, 1, -1)$ in $\mathfrak{F}$ ein. Durch die Lage von R und C auf $\mathfrak{F}$ wird zugleich gesichert, daß $0, 1, 1$ nicht ins Innere von $\mathfrak{R}$ fällt.

Nun möge noch $F(\pm 1, \pm 1, 0) > 1$ geblieben sein. Die Strecken $S'A$ und RB sind beide parallel und gleich lang mit OC . Wir führen jetzt jene Parallelprojektion von $\mathfrak{R}$ in Richtung OC auf eine Ebene, die uns für den ganzen Körper $\mathfrak{R}$ die Figur $\mathfrak{P}$ ergab, speziell nur für alle solchen zu OC parallelen Sehnen von $\mathfrak{R}$ aus, welche an Länge $\geq OC$ sind. Dadurch erhalten wir in jener Ebene eine neue Figur $\mathfrak{Q}$ mit O als Mittelpunkt, die ganz in der früheren Figur $\mathfrak{P}$ enthalten ist und insbesondere die Punkte $\mathfrak{A}$ und $\mathfrak{B}$ aufnimmt (Fig. 2). Dieser Figur $\mathfrak{Q}$ kommt, weil $\mathfrak{R}$ ein konvexer Körper ist, ebenfalls die Eigenschaft zu, mit irgend zwei Punkten stets die ganze sie verbindende Strecke zu enthalten, und stellt mithin $\mathfrak{Q}$ wieder ein konvexes Oval vor.

Liegt $\mathfrak{B}$ auf dem Rande von $\mathfrak{Q}$, so bedenken wir, daß solchen Punkten des Randes von $\mathfrak{Q}$, welche Häufungsstellen von nicht zu $\mathfrak{Q}$ zählenden Punkten aus $\mathfrak{P}$ sind, mit OC parallele Sehnen in $\mathfrak{R}$ *genau* von der Länge OC entsprechen und daß andererseits solchen Punkten des Randes von $\mathfrak{Q}$, die auch zum Rande von $\mathfrak{P}$ gehören, mit OC parallele Geraden entsprechen, die *nur die Begrenzung* von $\mathfrak{R}$ treffen. Infolgedessen können wir irgendeine kontinuierliche Veränderung von $\mathfrak{B}$ auf dem Rande von $\mathfrak{Q}$, während A und C festbleiben sollen, immer so bewerkstelligen, daß dabei sowohl B wie der damit in bestimmter Weise verbundene Punkt R fortwährend auf der Begrenzung von $\mathfrak{R}$ verbleiben. Wir können nun im allgemeinen $\mathfrak{B}$ auf dem Rande von $\mathfrak{Q}$ kontinuierlich der Geraden $O\mathfrak{A}$ nähern und damit Δ verringern. Nur wenn $\mathfrak{B}$ innerhalb einer diesem

Rande angehörigen, zu $O\mathfrak{A}$ parallelen Strecke liegt, müssen wir zuvor $\mathfrak{B}$ nach einem Endpunkte dieser Strecke führen, eine Operation, während der Δ sich nicht ändert. Wir bemerken, daß bei der letzteren Sachlage die Fläche $\mathfrak{F}$ sicher eine geradlinige Strecke aufzuweisen hat, nämlich, falls $\mathfrak{B}$ zugleich dem Rande von $\mathfrak{P}$ angehört, die Strecke RB , anderenfalls aber in einer Stützebene an $\mathfrak{R}$ durch B die ganze Strecke, als deren Projektion hier jene erwähnte Strecke auf dem Rande von $\mathfrak{Q}$ erscheint.

Sollte $\mathfrak{B}$ im Inneren von $\mathfrak{Q}$ liegen, so würden durch B und R jedenfalls zwei verschiedene Stützebenen an $\mathfrak{R}$ gehen; für diese müßten dann in gewissen Umgebungen von B und R die in Richtung OC gemessenen Abstände durchweg $\geq OC$ sein, also wären die Ebenen parallel, und könnten wir B in seiner Ebene so verändern, daß $\mathfrak{B}$ geradlinig auf O zuschreitet. Freilich würde unsere Folgerung hier mit der anderen Annahme in Widerspruch kommen, wonach C auf $\mathfrak{F}$ selbst liegen sollte.

Nach dieser Ausführung können wir nun, ohne daß Δ wächst, zu einem Zustande fortschreiten, wobei schließlich noch in einer der Ungleichungen $F(\pm 1, \pm 1, 0) \geq 1$ das Zeichen $=$ eintritt. Würden wir $F(1, 1, 0) = 1$ erhalten, so hätten wir in $-1, 0, 1; 0, -1, 1; 1, 1, 0$ drei Gitterpunkte auf $\mathfrak{F}$, die ein Gitteroktaeder zweiter Art bestimmen. Da wir solches vorläufig ausgeschlossen haben, so bleibt nur die Annahme übrig, daß noch der Gitterpunkt $T(1, -1, 0)$ auf $\mathfrak{F}$ auftritt.

18. Wenn die sechs Gitterpunkte A, B, C, R, S, T sämtlich auf $\mathfrak{F}$ liegen, so ersehen wir aus

$$F(1, 1, 2) + F(-1, 0, 1) + F(0, -1, 1) \geq 4F(0, 0, 1),$$

$$F(1, -1, 2) + F(-1, 1, 0) \geq 2F(0, 0, 1),$$

$$F(-1, -1, 2) + F(1, -1, 0) \geq 2F(0, -1, 1),$$

$$F(1, 1, 1) + F(-1, 0, 1) + F(0, -1, 1) \geq 3F(0, 0, 1)$$

usf., daß von den Ungleichungen (11), (12), (13) nur noch die Erfüllung der folgenden

$$F(-1, 1, 1) \geq 1, \quad F(1, -1, 1) \geq 1, \quad F(1, 1, -1) \geq 1$$

besonders zu fordern ist. Würde $F(1, 1, -1) = 1$ sein, so hätten wir in $0, 0, 1; 1, -1, 0; 1, 1, -1$ drei Gitterpunkte auf $\mathfrak{F}$, die ein Gitteroktaeder zweiter Art bestimmen. Gegenwärtig haben wir danach in diesen drei Ungleichungen das Zeichen $>$ zu verlangen.

19. Wir berücksichtigen jetzt die in 17. zunächst ausgeschaltete Möglichkeit, daß in einer der Ungleichungen $F(\pm 1, \pm 1, \pm 1) \geq 1$ das Gleichheitszeichen eintritt, schließen aber immer noch aus, daß auf $\mathfrak{F}$ drei Gitterpunkte auftreten, die ein Gitteroktaeder zweiter Art bestimmen.

Nehmen wir z. B. an, daß neben A, B, C noch der Gitterpunkt $D(1, 1, 1)$ auf $\mathfrak{F}$ liege, so sind wegen

$$F(-1, -1, 1) + F(1, 1, 1) \geq 2F(0, 0, 1)$$

usf. alle Ungleichungen (12) sichergestellt. Die Strecke AD ist parallel und gleich $B'C$. Wir verfahren unter Zuhilfenahme einer Parallelprojektion von $\mathfrak{R}$ in dieser Richtung wesentlich nach der in 17. dargelegten Methode zur Verringerung von Δ , und wir dürfen nun ferner in einer der Un-

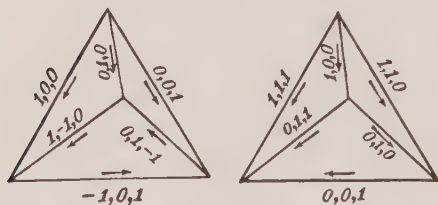


Fig. 3.

gleichungen (11) das Gleichheitszeichen als gültig annehmen. Würde dabei etwa $F(1, -1, 0) = 1$ eintreten, so hätten wir in $0, 0, -1$; $1, 1, 1$; $1, -1, 0$ drei Gitterpunkte auf $\mathfrak{F}$, die ein Gitteroktaeder zweiter Art bestimmen, was nicht sein sollte.

Wir nehmen nun etwa $N(1, 1, 0)$ auf $\mathfrak{F}$ gelegen an. Alsdann sind die Strecken $ON, CD, A'B$ parallel und gleich lang; wir verwenden eine Parallelprojektion von $\mathfrak{R}$ in Richtung ON und dürfen endlich noch etwa $L(0, 1, 1)$ auf $\mathfrak{F}$ gelegen annehmen. Die beistehende Figur zeigt uns, daß durch die Substitution

$$X^* = Z, \quad Y^* = X - Y, \quad Z^* = Y - Z$$

dieser Fall auf den schon in 17. behandelten zurückführt.

20. Wir nehmen jetzt zweitens an, daß das Okt (ABC) ein *Gitteroktaeder zweiter Art* ist, und haben alsdann neben den drei Gleichungen (10) die Ungleichungen (14)

$$F\left(\pm \frac{1}{2}, \pm \frac{1}{2}, \pm \frac{1}{2}\right) \geq 1$$

vorauszusetzen.

Zunächst möge in keiner dieser Ungleichungen das Zeichen $=$ gelten. Wir verfahren wesentlich nach der in 17. entwickelten Methode. Wir verwenden eine Parallelprojektion von $\mathfrak{R}$ in Richtung OC und können $\alpha_1, \dots, \gamma_3$ so variieren, daß Δ abnimmt oder konstant bleibt und schließlich in einer dieser Ungleichungen das Gleichheitszeichen eintritt. Es möge etwa der Gitterpunkt $N\left(\frac{1}{2}, \frac{1}{2}, -\frac{1}{2}\right)$ und sein Gegenpunkt N' auf

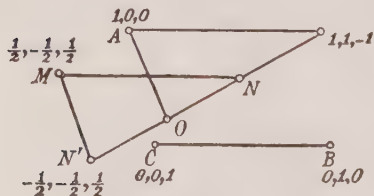


Fig. 4.

die Fläche $\mathfrak{F}$ fallen; in den übrigen Ungleichungen (14) aber möge noch das Zeichen $>$ bleiben. Nun ist CB parallel der Ebene durch O, A und N (Fig. 4).

Wir verwenden eine Parallelprojektion von $\mathfrak{R}$ in Richtung CB , halten A und die Strecke CB nach Richtung und Länge und

damit auch den Punkt N fest und können durch Variation von B und so, daß Δ abnimmt oder sich nicht ändert, einen Zustand erzielen, wobei eine

neue der Ungleichungen (14) als Gleichung erfüllt ist; es trete etwa der Gitterpunkt $L\left(-\frac{1}{2}, \frac{1}{2}, \frac{1}{2}\right)$ in die Fläche $\mathfrak{F}$ ein. — Dieser letzte Prozeß wäre nun genau so anzuwenden gewesen, wenn die Fläche $\mathfrak{F}$ neben N auch bereits $M\left(\frac{1}{2}, -\frac{1}{2}, \frac{1}{2}\right)$ enthalten hätte, da die Strecke MN parallel und gleich CB ist. Daraus leuchtet sofort ein, daß wir überhaupt auch einen Zustand erreichen können, wobei drei von den vier Paaren von Gitterpunkten $\pm\frac{1}{2}, \pm\frac{1}{2}, \pm\frac{1}{2}$, etwa die drei Paare L, L', M, M', N, N' , auf $\mathfrak{F}$ fallen.

Vermöge der Substitution

$$X^* = Y + Z, \quad Y^* = X + Z, \quad Z^* = X + Y$$

erhalten nun A, B, C, L, M, N und der Gitterpunkt $\frac{1}{2}, \frac{1}{2}, \frac{1}{2}$ die neuen Koordinaten

$$X^*, Y^*, Z^* = 0, 1, 1; 1, 0, 1; 1, 1, 0; 1, 0, 0; 0, 1, 0; 0, 0, 1; 1, 1, 1.$$

Das ursprüngliche Zahlengitter ($\mathfrak{G}$) wird mit dem Zahlengitter in X^*, Y^*, Z^* identisch; und damit kein Gitterpunkt außer O im Inneren von $\mathfrak{R}$ liegt, bleibt nur noch die eine Bedingung zu erfüllen, daß der Punkt $X^*, Y^*, Z^* = 1, 1, 1$ nicht ins Innere von $\mathfrak{R}$ fällt.

21. Wir sind durch diese Überlegungen zu folgendem Resultate gelangt:

Um für einen gegebenen konvexen Körper K das Minimum (oder die verschiedenen existierenden Minima) von Δ zu finden, genügt es, solche Anordnungen des Gitters ($\mathfrak{G}$) in Betracht zu ziehen, wobei von diesem Gitter entweder (I) auf die Begrenzung von $\mathfrak{R} = K + K'$ die Punkte

$$1, 0, 0; 0, 1, 0; 0, 0, 1; 0, 1, -1; -1, 0, 1; 1, -1, 0$$

und die Punkte $-1, 1, 1; 1, -1, 1; 1, 1, -1$ außerhalb $\mathfrak{R}$ fallen, oder (II) auf die Begrenzung von $\mathfrak{R}$ die Punkte

$$1, 0, 0; 0, 1, 0; 0, 0, 1; 0, 1, 1; 1, 0, 1; 1, 1, 0$$

fallen und der Punkt $1, 1, 1$ außerhalb $\mathfrak{R}$ liegt,

oder (III) auf die Begrenzung von $\mathfrak{R}$ die Punkte

$1, 0, 0; 0, 1, 0; 0, 0, 1; 0, 1, 1; 1, 0, 1; 1, 1, 0$ und $1, 1, 1$ fallen.

Aus einer beliebigen Anordnung des Gitters ($\mathfrak{G}$) von dem hier bezeichneten Charakter, welche ein Minimum von Δ liefert, können unter Umständen, — jedoch nur in solchen Fällen, wo auf der Begrenzung von $\mathfrak{R}$ geradlinige Strecken vorkommen, — andere Anordnungen von ($\mathfrak{G}$), welche nicht jenen Charakter tragen, durch kontinuierliche Variation ohne Änderung des Wertes von Δ hervorgehen und diese würden dann in

gleicher Weise dichteste gitterförmige Lagerungen für den Grundkörper K bestimmen.

Die in (I) aufgeführten sechs Gitterpunkte mit ihren Gegenpunkten in bezug auf O bilden die Ecken eines Kubooktaeders, die in (III) ge-

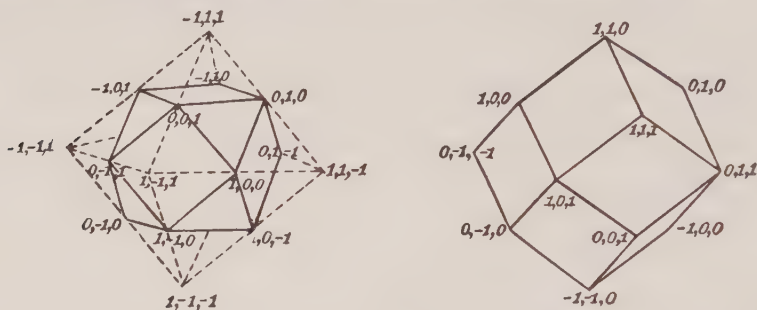


Fig. 5.

nannten sieben Gitterpunkte mit ihren Gegenpunkten in bezug auf O die Ecken eines Rhombendodekaeders (Fig. 5).

§ 7. Arithmetische Äquivalenz bei positiven ternären quadratischen Formen.

22. Ehe wir in der Behandlung unseres allgemeinen Problems weitergehen, wollen wir das Beispiel der dichtesten Lagerung von Kugeln näher ins Auge fassen.

Es sei K eine Kugel vom Radius $\frac{1}{2}$, für den Körper $\mathfrak{K} = K + K'$ also $\varphi(\xi, \eta, \zeta) = (\xi^2 + \eta^2 + \zeta^2)^{\frac{1}{2}}$, so haben wir in

$$\begin{aligned} \varphi^2 &= (\alpha_1 x + \alpha_2 y + \alpha_3 z)^2 + (\beta_1 x + \beta_2 y + \beta_3 z)^2 + (\gamma_1 x + \gamma_2 y + \gamma_3 z)^2 \\ &= e_{11}x^2 + 2e_{12}xy + 2e_{13}xz + e_{22}y^2 + 2e_{23}yz + e_{33}z^2 = h(x, y, z) \end{aligned}$$

eine positive ternäre quadratische Form mit einer Determinante $D = \Delta^2$.

Um die dichteste gitterförmige Lagerung für Kugeln vom Radius $\frac{1}{2}$ zu finden, haben wir nach 21. nur diese zwei Annahmen zu diskutieren:

(II) Die Fläche $h(x, y, z) = 1$ enthält sechs Gitterpunkte

$$x, y, z = 1, 0, 0; 0, 1, 0; 0, 0, 1; 0, -1, -1; 1, 0, -1; 1, 1, 0;$$

dann müßte

$$h(x, y, z) = x^2 - xy - xz + y^2 - yz + z^2 = \left(x - \frac{1}{2}y - \frac{1}{2}z\right)^2 + \frac{3}{4}(y - z)^2$$

sein, und jene Fläche wäre eine Zylinderfläche, nicht eine Kugel.

(I) Die Fläche $h(x, y, z) = 1$ enthält die sechs Gitterpunkte:

$$x, y, z = 1, 0, 0; 0, 1, 0; 0, 0, 1; 0, 1, -1; -1, 0, 1; 1, -1, 0;$$

dann haben wir

$$h(x, y, z) = x^2 + xy + xz + y^2 + yz + z^2$$

und wird $D = \Delta^2 = \frac{1}{2}$. Demnach gelangen wir zu dem Resultate:

Im Falle der dichtesten gitterförmigen Lagerung von gleichen Kugeln verhält sich der von den Kugeln erfüllte Raum zum ganzen unendlichen Raume wie $\frac{\pi}{6} : \frac{1}{\sqrt{2}}$. Die betreffende Lagerung ist dadurch charakterisiert, daß eine jede Kugel in den zwölf Ecken eines Kubooktaeders an andere Kugeln stößt.

Für die dichteste Lagerung von Ellipsoiden gilt offenbar genau der nämliche Satz.

Das entsprechende Problem für zwei Dimensionen hängt analog mit der Transformation von $\xi^2 + \eta^2$ in $x^2 + xy + y^2$ zusammen, und kann man in einer Ebene durch lauter gleiche Kreisflächen, die nicht ineinander eindringen und deren Mittelpunkte ein parallelogrammatisches Punktsystem bilden, im Maximum den $\frac{\pi}{4} : \frac{\sqrt{3}^{\text{ten}}}{2}$ Teil der ganzen unendlichen Ebene ausfüllen.

23. Ich will jetzt zeigen, daß die *Theorie der arithmetischen Äquivalenz der positiven ternären quadratischen Formen*, deren Hauptsätze von Seeber und Gauß aufgestellt sind und in der Folge manche andere Ableitung gefunden haben*), vollständig und durch einfache Überlegungen allein aus der eben bewiesenen Tatsache über die dichteste Lagerung von Kugeln erschlossen werden kann.

Es sei

$$e_{11}x^2 + 2e_{12}xy + \dots + e_{33}z^2 = h(x, y, z) = (f(x, y, z))^2$$

eine beliebige positive ternäre quadratische Form mit der positiven Determinante D . Wir setzen die Form h irgendwie in die Gestalt

$h = (\alpha_1x + \alpha_2y + \alpha_3z)^2 + (\beta_1x + \beta_2y + \beta_3z)^2 + (\gamma_1x + \gamma_2y + \gamma_3z)^2$
 $= \xi^2 + \eta^2 + \zeta^2$ und deuten ξ, η, ζ als rechtwinklige Koordinaten im Raume. Der Ausdruck $\sqrt{h} = f(x, y, z)$ stellt alsdann die Entfernung des variablen Punktes P mit den Bestimmungsstücken x, y, z vom Nullpunkte O dar.

24. Wir betrachten das Gitter ($\mathfrak{G}$) der Punkte mit ganzzahligen Werten x, y, z und bestimmen in diesem Gitter einen ersten vom Null-

*) Seeber, Untersuchungen über die Eigenschaften der positiven ternären quadratischen Formen, Freiburg i. Br., 1831. — Gauß, Göttingische gelehrte Anzeigen, Jahrg. 1831 (Werke, Bd. II, S. 188). — Dirichlet, Crelles Journal, Bd. 40, S. 209 (Werke, Bd. II, S. 27). — Hermite, Crelles Journal, Bd. 40, S. 173; Bd. 79, S. 17 (Oeuvres T. I, p. 94 und T. III). — Selling, Crelles Journal, Bd. 77, S. 143. — Korkine u. Zolotareff, Mathematische Annalen, Bd. 6, S. 366.

punkte verschiedenen Gitterpunkt A derart, daß die Entfernung OA so klein als möglich ausfällt, hernach einen zweiten Gitterpunkt B außerhalb der Geraden durch O und A , so daß die Entfernung OB so klein als möglich wird, endlich einen dritten Gitterpunkt C außerhalb der Ebene durch O , A und B , so daß die Entfernung OC so klein als möglich wird. Es seien $l_1, m_1, n_1; l_2, m_2, n_2; l_3, m_3, n_3$ die x, y, z -Werte von A, B, C und f_1, f_2, f_3 die Entfernungen dieser Punkte von O .

Wir setzen

(15) $x = l_1 X + l_2 Y + l_3 Z$, $y = m_1 X + m_2 Y + m_3 Z$, $z = n_1 X + n_2 Y + n_3 Z$, und es sei d die Determinante dieser Ausdrücke. Die Form h gehe durch diese lineare Substitution in

$$E_{11} X^2 + 2E_{12} XY + \dots + E_{33} Z^2 = H(X, Y, Z)$$

über; dabei wird $E_{11} = f_1^2$, $E_{22} = f_2^2$, $E_{33} = f_3^2$. Wir bringen H in die Gestalt

$$H = (A_1 X + A_2 Y + A_3 Z)^2 + (B_2 Y + B_3 Z)^2 + (\Gamma_3 Z)^2,$$

so daß $A_1, B_2, \Gamma_3 > 0$ sind, und setzen

$$\Phi = \left(\frac{A_1 X + A_2 Y + A_3 Z}{f_1} \right)^2 + \left(\frac{B_2 Y + B_3 Z}{f_2} \right)^2 + \left(\frac{\Gamma_3 Z}{f_3} \right)^2.$$

25. Die Fläche $\Phi = 1$ stellt im Raume der rechtwinkligen Koordinaten ξ, η, ζ ein Ellipsoid vor mit einer Hauptachse in der Linie OA , einer dazu senkrechten Hauptachse in der Ebene OAB und einer dritten auf dieser Ebene senkrechten Hauptachse, bzw. von den Längen f_1, f_2, f_3 . Für jeden von O verschiedenen Gitterpunkt in der Geraden OA ($Y = 0, Z = 0$) ist $H \geq f_1^2$, $\Phi \geq 1$, für jeden Gitterpunkt außerhalb dieser Geraden in der Ebene OAB ($Z = 0$) ist $H \geq f_2^2$, $\Phi \geq 1$, für jeden Gitterpunkt außerhalb der Ebene OAB ist $H \geq f_3^2$, $\Phi \geq 1$. Danach liegt kein Gitterpunkt x, y, z außer dem Nullpunkte im Inneren des Ellipsoids $\Phi \leq 1$.

Konstruieren wir nun den Körper $\Phi \leq \frac{1}{2}$ und weiter alle Körper, die aus diesem durch die Translationen vom Nullpunkte nach den einzelnen Gitterpunkten x, y, z entstehen, so werden diese sämtlichen gitterförmig angeordneten Ellipsoide untereinander höchstens Punkte der Begrenzungen gemein haben, und ist daher nach 22. das Verhältnis aus dem Volumen von $\Phi \leq \frac{1}{2}$ und dem Volumen des Grundparallelepipeds des Gitters ($\mathfrak{G}$), also $\frac{\pi}{6} f_1 f_2 f_3 : \sqrt{D}$ sicher $\leq \frac{\pi}{6} : \frac{1}{\sqrt{2}}$, d. h. wir haben

$$(16) \quad f_1 f_2 f_3 \leq \sqrt{2D}, \quad E_{11} E_{22} E_{33} \leq 2D.$$

Andererseits berechnet sich die Determinante von H in x, y, z zu $\left(\frac{A_1 B_2 \Gamma_3}{d} \right)^2$, so daß

$$A_1 B_2 \Gamma_3 = |d| \sqrt{D}$$

folgt. Nun ist für die Punkte A, B, C bzw. $H = f_1^2, f_2^2, f_3^2$, so daß

$$A_1 = f_1, B_2 \leq f_2, \Gamma_3 \leq f_3$$

entsteht. Demnach gilt

$$f_1 f_2 f_3 \geq |d| \sqrt{D}.$$

Mit Berücksichtigung von (16) geht hieraus

$$(17) \quad |d| \leq \sqrt{2},$$

mithin $d = \pm 1$ hervor. Die Determinante der Substitution (15) ist also ± 1 , das Gitter in X, Y, Z ist identisch mit dem Gitter $(\mathfrak{G})$ in x, y, z , die Form H arithmetisch äquivalent der gegebenen Form h .

26. Nach der Art, wie die Punkte A, B, C mit ihren Entfernungen OA, OB, OC eingeführt wurden, bestehen für die Form H die sämtlichen Ungleichungen

$$(18) \quad 0 < E_{11} \leq E_{22} \leq E_{33},$$

$$(19) \quad \begin{array}{ccc} H(X, 0, 0) \geq E_{11}, & H(X, Y, 0) \geq E_{22}, & H(X, Y, Z) \geq E_{33}, \\ (X \neq 0) & (Y \neq 0) & (Z \neq 0) \end{array}$$

worin X, Y, Z beliebige ganze Zahlen sind. Umgekehrt leuchtet ein, daß mit diesen Ungleichungen, während die Substitution (15) ganzzahlige Koeffizienten und eine Determinante ± 1 hat, der hier in Frage kommende Charakter der Gitterpunkte A, B, C vollständig erschöpft wird.

Eine ternäre quadratische Form $H(X, Y, Z)$, welche diese sämtlichen Ungleichungen (18), (19) erfüllt, heißt *reduziert*.

Analog heißt eine binäre quadratische Form

$$H(X, Y) = E_{11}X^2 + 2E_{12}XY + E_{22}Y^2$$

reduziert, wenn sie allen Ungleichungen

$$0 < E_{11} \leq E_{22}, \quad \begin{array}{cc} H(X, 0) \geq E_{11}, & H(X, Y) \geq E_{22} \\ (X \neq 0) & (Y \neq 0) \end{array}$$

genügt, worin X, Y beliebige ganze Zahlen sind. Es leuchtet ein, daß, wenn $H(X, Y, Z)$ eine reduzierte ternäre Form ist, die drei daraus durch Nullsetzen je einer Variable entstehenden binären Formen $H(X, Y, 0)$, $H(X, 0, Z)$, $H(0, Y, Z)$ ihrerseits notwendig reduziert sind.

Aus der Menge der Ungleichungen (19) greifen wir insbesondere die folgenden heraus:

$$(20) \quad H(\pm 1, 1, 0) \geq E_{22}, \quad H(\pm 1, 0, 1) \geq E_{33}, \quad H(0, \pm 1, 1) \geq E_{33},$$

$$\text{d. i.} \quad E_{11} \geq 2|E_{12}|, \quad E_{11} \geq 2|E_{13}|, \quad E_{22} \geq 2|E_{23}|,$$

und ferner

$$(21) \quad H(\pm 1, \pm 1, 1) \geq E_{33};$$

diese letzteren Ungleichungen (21) sind infolge der Ungleichungen (20) von selbst erfüllt, wenn die Größen E_{12}, E_{13}, E_{23} alle drei positiv oder

eine positiv und zwei negativ sind oder darunter wenigstens eine verschwindende auftritt; wenn aber diese Größen alle drei negativ oder zwei von ihnen positiv, eine negativ sind, so liefert (21) die eine neue Bedingung

$$E_{11} + E_{22} \geq 2 |E_{12}| + 2 |E_{13}| + 2 |E_{23}|.$$

Wir wollen jetzt den Nachweis erbringen, daß die speziellen, in endlicher Anzahl vorhandenen Ungleichungen (18), (20), (21) die sämtlichen unendlich vielen Ungleichungen (19) nach sich ziehen und also den Charakter von H als reduzierte Form bereits völlig bestimmen.

27. Wir betrachten zu dem Ende die Mannigfaltigkeit der sechs unabhängigen Variablen $E_{11}, E_{12}, \dots, E_{33}$ und in dieser denjenigen Bereich ($\mathfrak{S}$), der durch die sämtlichen Ungleichungen (18), (19) für diese Variablen definiert ist, wobei wir noch die Bedingung $E_{11} > 0$ durch $E_{11} \geq 0$ ersetzen wollen. Jedem Punkte $E_{11}, E_{12}, \dots, E_{33}$ in dieser reduzierten Kammer ($\mathfrak{S}$) entspricht eine niemals negative ternäre Form H . Da die Ungleichungen (18), (19) linear und homogen in den Variablen $E_{11}, E_{12}, \dots, E_{33}$ sind, so besitzt ($\mathfrak{S}$) die Eigenschaft, mit irgend zwei Punkten (Formen) $H^{(0)}, H^{(1)}$ stets die ganze sie verbindende Strecke, d. h. die Koeffizientensysteme der Formen $(1-t)H^{(0)} + tH^{(1)}$ für alle Parameterwerte $t \geq 0$ und ≤ 1 zu enthalten, stellt also einen konvexen Körper in der Mannigfaltigkeit der $E_{11}, E_{12}, \dots, E_{33}$ dar. Nach den allgemeinen Grundsätzen über die Begrenzung eines konvexen Körpers*) genügt es nun zur Definition des Bereichs ($\mathfrak{S}$), von den Ungleichungen (18), (19) nur jede solche ausdrücklich zu fordern, für welche im Bereiche ($\mathfrak{S}$) ein Punkt gefunden werden kann, der die betreffende Ungleichung mit dem Zeichen $=$, alle davon verschiedenen der Ungleichungen (18), (19) aber mit dem Zeichen $> (<)$ erfüllt.

Jeder Punkt H in ($\mathfrak{S}$), der nicht einer wesentlich positiven Form entspricht, ist zufolge der angegebenen Eigenschaft von ($\mathfrak{S}$) gewiß Häufungsstelle von wesentlich positiven reduzierten Formen und aus der Ungleichung (16) für diese letzteren Formen geht dann die nämliche Ungleichung auch für H hervor; daher ist für H dann notwendig $D = 0$, $E_{11} = 0$. Die letztere Relation hat wegen (20) noch $E_{12} = 0$, $E_{13} = 0$ zur Folge, und ist die Ungleichung $E_{11} \geq 0$ hiernach bei der Definition von ($\mathfrak{S}$) entbehrlich. Jeder Punkt von ($\mathfrak{S}$) aber, für welchen $E_{11} > 0$ ist, liefert gewiß eine wesentlich positive Form.

28. Greifen wir nunmehr eine in dem eben besprochenen Sinne notwendige der Ungleichungen (18), (19) heraus; es sei dieses etwa eine Ungleichung

*) Geometrie der Zahlen, Leipzig, 1896.

$$H(a, b, c) \geq E_{33},$$

wo a, b, c ganze Zahlen sind und $c \neq 0$ ist. Dann gibt es also im Bereich $(\mathfrak{H})$ irgendeine wesentlich positive Form $H = (E_{11}, E_{12}, \dots, E_{33})$, deren Koeffizienten diese Ungleichung mit dem Zeichen $=$, alle davon verschiedenen der Ungleichungen (18), (19) aber mit dem Zeichen $>$ erfüllen.

Bezeichnen wir, wie in 24., mit $OABC$ das Grundparallelepiped des Gitters in X, Y, Z , so liegt wegen $c \neq 0$ der Gitterpunkt G ($X = a$, $Y = b$, $Z = c$) außerhalb der Ebene OAB und bietet hier dieselbe Entfernung $\sqrt{H}$ von O wie C dar. Wir könnten daher bei der Aufsuchung einer beliebigen mit H äquivalenten reduzierten Form den Punkt G an die Stelle von C treten lassen, und nach (17) muß alsdann die Determinante aus den Koordinaten von A, B, G notwendig ± 1 sein, d. h. wir haben $c = \pm 1$.

Aus der am Schlusse von 22. hinzugefügten Bemerkung über die dichteste gitterförmige Lagerung von Kreisflächen in der Ebene schließen wir analog, daß zur Charakterisierung einer reduzierten binären Form

$$H(X, Y) = E_{11}X^2 + 2E_{12}XY + E_{22}Y^2$$

jedenfalls die Ungleichungen

$$0 \leq E_{11} \leq E_{22}, \quad H(X, 1) \geq E_{22}$$

ausreichen.

Wir ersehen daraus zunächst, daß wir bei der Definition einer reduzierten ternären Form von den Ungleichungen (19) gewiß nur die folgenden nötig haben

$$H(X, 1, 0) \geq E_{22}, \quad H(X, Y, 1) \geq E_{33}.$$

Nehmen wir jetzt an, es sei $b \neq 0$. Da in den hier aufgeführten Ungleichungen die Größe E_{33} herausfällt, so können wir in der Form $H(X, Y, Z)$ an Stelle des Koeffizienten E_{33} den Wert $E_{33}^* = E_{22}$ setzen und die dadurch entstehende neue Form H^* wird ebenfalls reduziert sein. Für diese Form H^* ist nun

$$H^*(a, b, c) = E_{22}^*.$$

Da wir $b \neq 0$ haben, so können wir bei der Aufsuchung einer beliebigen mit H^* äquivalenten reduzierten Form den Punkt G an die Stelle von B treten lassen, während wir A und C unverändert beibehalten, und es muß nunmehr die Determinante aus den Koordinaten von A, G, C notwendig ± 1 , also $b = \pm 1$ sein.

Der analoge Schluß für binäre Formen zeigt, daß eine reduzierte binäre Form $H(X, Y) = (E_{11}, E_{12}, E_{22})$ völlig durch die Ungleichungen

$$0 \leq E_{11} \leq E_{22}, \quad H(\pm 1, 1) \geq E_{22}$$

charakterisiert ist.

Mit Rücksicht hierauf brauchen wir nunmehr für eine reduzierte ternäre Form $H(X, Y, Z)$ von den Ungleichungen (19) nur noch die folgenden in Betracht zu ziehen:

$$H(\pm 1, 1, 0) \geq E_{22}, \quad H(\pm 1, 0, 1) \geq E_{33}, \quad H(X, \pm 1, 1) \geq E_{33}.$$

Nehmen wir nun an, es sei z. B. $c = 1$, $b = -1$ und $a \neq 0$. Wir ersetzen in H^* die Koeffizienten E_{22}^* , E_{23}^* , E_{33}^* durch $E_{22}^* - \varepsilon$, $E_{23}^* - \frac{1}{2}\varepsilon$, $E_{33}^* - \varepsilon$, wobei $\varepsilon > 0$ sei. Die neu entstehende Form H^{**} erfüllt in völlig unveränderter Weise die Beziehungen

$$E_{22}^{**} = E_{33}^{**}, \quad H^{**}(a, b, c) = E_{22}^{**},$$

$$H^{**}(\pm 1, 1, 0) \geq E_{22}^{**}, \quad H^{**}(\pm 1, 0, 1) \geq E_{33}^{**}, \quad H^{**}(X, -1, 1) \geq E_{33}^{**}.$$

Wir setzen $\varepsilon \leq E_{22}^* - E_{11}^*$ voraus, und es bleibt auch $E_{11}^{**} \leq E_{22}^{**}$. Sollte nun, wenn wir ε von 0 an bis zu der hier genannten Grenze kontinuierlich wachsen lassen, schließlich einmal in irgendeiner Ungleichung $H^{**}(X, 1, 1) \geq E_{33}^{**}$ das Gleichheitszeichen eintreten, so könnten wir für die betreffende immer noch reduzierte Form H^{**} die Punkte $1, 0, 0$; $X, 1, 1$; $a, -1, 1$ an die Stelle von A, B, C treten lassen; die Determinante aus den Koordinaten dieser Punkte aber wäre $= 2$ im Widerspruch mit der Bedingung (17). Also muß H^{**} auch noch bis zu $\varepsilon = E_{22}^* - E_{11}^*$ reduziert bleiben. Bei diesem Werte von ε haben wir nun $E_{11}^{**} = E_{22}^{**}$. Also kann für H^{**} jetzt der Punkt G an die Stelle von A treten, während B und C ihre Rollen beibehalten, und muß endlich die Determinante aus den Koordinaten von G, B, C , also $a = \pm 1$ sein. Damit ist in der Tat die in 26. aufgestellte Behauptung erwiesen.

29. Die erlangten Sätze sprechen wir in folgender Weise aus:

Zu einer beliebig gegebenen ternären positiven quadratischen Form $h(x, y, z)$ kann man immer eine ganzzahlige lineare Substitution mit der Determinante ± 1 bestimmen, durch welche $h(x, y, z)$ in eine Form

$$H(X, Y, Z) = E_{11}X^2 + 2E_{12}XY + \cdots + E_{33}Z^2$$

übergeht, die den Ungleichungen

$$0 < E_{11} \leq E_{22} \leq E_{33}, \quad 2|E_{12}| \leq E_{11}, \quad 2|E_{13}| \leq E_{11}, \quad 2|E_{23}| \leq E_{22}$$

und zudem, wenn $E_{12}E_{13}E_{23} < 0$ ist, noch der Ungleichung

$$2|E_{12}| + 2|E_{13}| + 2|E_{23}| \leq E_{11} + E_{22}$$

genügt.

Die Form H behält diesen Charakter, wenn man darin noch irgendwie X durch $\pm X$, Y durch $\pm Y$, Z durch $\pm Z$ ersetzt. Abgesehen von diesen möglichen Abänderungen ist jene Substitution und diese reduzierte mit h äquivalente Form H völlig bestimmt, wofern in keiner der genannten Ungleichungen das Gleichheitszeichen statthat. In jedem Falle aber sind die

Werte E_{11}, E_{22}, E_{33} eindeutig bestimmt. Dabei folgen aus den angenommenen Ungleichungen die sämtlichen Ungleichungen:

$$\begin{aligned} H(X, Y, Z) &\geq E_{11}, & H(X, Y, Z) &\geq E_{22}, & H(X, Y, Z) &\geq E_{33} \\ (X, Y, Z \neq 0, 0, 0) & & (Y, Z \neq 0, 0) & & (Z \neq 0) \end{aligned}$$

für ganzzahlige X, Y, Z .

Endlich gilt dabei die Beziehung (Theorem von Gauß):

$$E_{11} E_{22} E_{33} \leq 2D.$$

In dieser Ungleichung tritt bei gegebenem D das Gleichheitszeichen nur ein, wenn

$$\begin{aligned} H &= \sqrt[3]{2D}(X^2 + Y^2 + Z^2 + XY + XZ + YZ), \text{ bzw.} \\ &= \sqrt[3]{2D}(X^2 + Y^2 + Z^2 - XY - YZ) \end{aligned}$$

ist bzw. daraus durch Permutation und Vorzeichenänderung der Variablen hervorgeht, welche speziellen reduzierten Formen sämtlich untereinander äquivalent sind (vgl. auch Fig. 3 in 19).

§ 8. Die weiteren Bedingungen für eine dichteste Lagerung.

30. Wir nehmen jetzt die Behandlung des Problems der dichtesten Lagerung für einen beliebigen konvexen Grundkörper K wieder auf. Die Größe Δ , welche wir für K unter Erfüllung gewisser Ungleichungen zu einem Minimum zu machen suchen, ist eine Funktion der neun Variablen $\alpha_1, \alpha_2, \dots, \alpha_9$, die uns zur Festlegung eines Gitters dienen. Bei einem Minimum von Δ können wir nach 21. gewisse 6 bzw. 7 Gleichungen als erfüllt voraussetzen. Es fehlen uns daher noch 3 bzw. 2 weitere Gleichungen zur Charakterisierung eines Minimums von Δ , die wir jetzt ermitteln wollen.

31. Wir verfolgen zuerst die Umstände des Falles (I) in 21.:

Es sollen die Punkte

$$\begin{aligned} \mathfrak{L}, \quad \mathfrak{M}, \quad \mathfrak{N}, \quad \mathfrak{R}, \quad \mathfrak{S}, \quad \mathfrak{T} \\ 1, 0, 0; 0, 1, 0; 0, 0, 1; 0, 1, -1; -1, 0, 1; 1, -1, 0 \end{aligned}$$

auf der Fläche $\mathfrak{F}$, der Begrenzung von $\mathfrak{R}$, und die Punkte $-1, 1, 1; 1, -1, 1; 1, 1, -1$ außerhalb $\mathfrak{R}$ liegen. Ferner setzen wir voraus, daß nicht auf $\mathfrak{F}$ drei Gitterpunkte vorhanden sind, die ein Gitteroktaeder zweiter Art bestimmen.

Können nun überhaupt außer den genannten 6 Gitterpunkten und ihren 6 Gegenpunkten noch andere Gitterpunkte auf der Fläche $\mathfrak{F}$ auftreten? Nach 15. können neben den Punkten $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}$ auf der Fläche $\mathfrak{F}$ überhaupt nur solche Gitterpunkte da sein, deren Koordinaten abgesehen

von der Reihenfolge und den Vorzeichen eines der folgenden Systeme ergeben:

$$0, 1, 1; 1, 1, 1; 0, 1, 2; 1, 1, 2; 1, 2, 2; \\ 1, 1, 3; 1, 2, 3; 2, 2, 3; 1, 2, 4; 2, 3, 4.$$

Nun leiten wir aus der allgemeinen Funktionalungleichung (8) für einen konvexen Körper insbesondere die folgenden Ungleichungen ab:

$$F\left(-\frac{2}{2}, 3, \pm 4\right) + F(-1, 1, 0) + F(-1, 0, 0) \geq 4F\left(\frac{0}{-1}, 1, \pm 1\right), \\ F\left(-1, -\frac{2}{2}, \pm 4\right) + F(1, -1, 0) + F(0, -1, 0) \geq 4F\left(0, \frac{0}{-1}, \pm 1\right), \\ F(a, b, c) + aF(-1, 0, 1) + bF(0, -1, 1) \geq (a + b + c)F(0, 0, 1), \\ (a, b, c = 2, 2, 3; 1, 2, 3; 1, 1, 3; 1, 2, 2; 1, 1, 2), \\ F\left(-\frac{2}{-1}, -2, 3\right) + F(1, -1, 0) \geq 3F\left(\frac{1}{0}, -1, 1\right), \\ F(-2, -2, 3) + F(-1, 0, 0) + F(0, -1, 0) \geq 3F(-1, -1, 1), \\ F(1, -2, 3) + F(1, -1, 0) + F(1, 0, 0) \geq 3F(1, -1, 1), \\ F(-1, 2, 3) + 2F(0, -1, 1) + F(1, 0, 0) \geq 5F(0, 0, 1), \\ F(-1, 1, 3) + F(1, -1, 0) \geq 3F(0, 0, 1), \\ F(-1, 2, 2) + F(-1, 0, 0) \geq 2F(-1, 1, 1), \\ F(1, -2, 2) + F(1, 0, 0) \geq 2F(1, -1, 1), \\ F(1, -1, 2) + F(1, -1, 0) \geq 2F(1, -1, 1), \\ F(0, 1, 2) + F(0, -1, 1) \geq 3F(0, 0, 1);$$

auf Grund dieser Ungleichungen sowie der weiteren, die daraus durch Permutation der Variablen hervorgehen, scheiden von den genannten Systemen sofort eine Reihe als außerhalb des Körpers $\mathfrak{K}$ gelegen aus. Es bleibt die Lage auf $\mathfrak{F}$ nur noch für diejenigen Gitterpunkte fraglich, deren Koordinaten abgesehen von der Reihenfolge eines der Systeme

$$-1, -1, 3; -1, -1, 2; 0, -1, 2; 1, 1, 1; 0, 1, 1$$

ergeben. Wenn aber einer dieser Gitterpunkte auf $\mathfrak{F}$ liegt, so erlangt man nach 11. entweder durch die Systeme $0, 0, -1; 1, -1, 0; -1, -1, 3$ oder $1, 0, 0; 0, 1, 0; -1, -1, 2$ oder $1, 0, 0; -1, 1, 0; 0, -1, 2$ oder $-1, 0, 0; 0, -1, 1; 1, 1, 1$ oder $1, -1, 0; 1, 0, -1; 0, 1, 1$ drei Gitterpunkte auf $\mathfrak{F}$, die ein Gitteroktaeder zweiter Art bestimmen. Wir haben demnach jetzt anzunehmen, daß die Fläche $\mathfrak{F}$ neben $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}, \mathfrak{R}, \mathfrak{S}, \mathfrak{T}$ und den Gegenpunkten keine weiteren Gitterpunkte aufweist.

Wir legen durch jeden der Punkte $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}, \mathfrak{R}, \mathfrak{S}, \mathfrak{T}$ eine Stützebene an $\mathfrak{K}$; die Gleichungen dieser Ebenen seien:

$$\begin{aligned}
 (22) \quad & L \equiv X + \lambda_2 Y + \lambda_3 Z = 1, \\
 & M \equiv \mu_1 X + Y + \mu_3 Z = 1, \\
 & N \equiv \nu_1 X + \nu_2 Y + Z = 1, \\
 & R \equiv (\varrho_2 - \varrho_1) X + \varrho_2 Y - (1 - \varrho_2) Z = 1, \\
 & S \equiv -(1 - \sigma_3) X + (\sigma_3 - \sigma_2) Y + \sigma_3 Z = 1, \\
 & T \equiv \tau_1 X - (1 - \tau_1) Y + (\tau_1 - \tau_3) Z = 1.
 \end{aligned}$$

Ein jeder Ausdruck L, M, N, R, S, T muß im ganzen Bereiche von $\mathfrak{R}$ im Intervalle ≥ -1 und ≤ 1 liegen. Die Form L wird für jene sechs Gitterpunkte bzw.

$$1, \lambda_2, \lambda_3, \lambda_2 - \lambda_3, -1 + \lambda_3, 1 - \lambda_2;$$

mithin sind λ_2, λ_3 beide ≥ 0 und ≤ 1 . Die Form R wird für die sechs Punkte bzw.

$$\varrho_2 - \varrho_1, \varrho_2, -1 + \varrho_2, 1, -1 + \varrho_1, -\varrho_1;$$

mithin sind ϱ_2, ϱ_1 beide ≥ 0 und ≤ 1 . Das nämliche gilt von $\mu_3, \mu_1; \nu_1, \nu_2; \sigma_3, \sigma_2; \tau_1, \tau_3$.

Wir denken uns nun die X, Y, Z -Koordinaten in ihrer augenblicklichen Bedeutung festgehalten und variieren dagegen das Gitter ($\mathfrak{G}$) kontinuierlich, indem wir die Punkte $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}$ an die Stellen verlegen, deren X, Y, Z -Koordinaten

$$1 + \varepsilon X_1, \varepsilon Y_1, \varepsilon Z_1; \varepsilon X_2, 1 + \varepsilon Y_2, \varepsilon Z_2; \varepsilon X_3, \varepsilon Y_3, 1 + \varepsilon Z_3$$

sind, unter ε einen positiven Parameter verstanden. Dabei sollen die Punkte $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}, \mathfrak{R}, \mathfrak{S}, \mathfrak{T}$ in den betreffenden Stützebenen oder auf deren dem Nullpunkte abgewandten Seiten bleiben, d. h. es sollen die Ungleichungen

$$L_1 \geq 0, M_2 \geq 0, N_3 \geq 0, R_2 - R_3 \geq 0, -S_1 + S_3 \geq 0, T_1 - T_2 \geq 0$$

gelten; wir bezeichnen hier die Werte der Formen $L, M, \dots, T$ für das System X_i, Y_i, Z_i durch Anhängen des Index i . Nach den vorausgeschickten Bemerkungen wird dabei jedenfalls kein Gitterpunkt ins Innere von $\mathfrak{R}$ eintreten, so lange ε eine gewisse Grenze nicht übersteigt.

Der Wert der Determinante Δ verändert sich bei dieser Variation von $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}$ aus $\Delta(0)$ in

$$\Delta(\varepsilon) = \Delta(0)(1 + \varepsilon(X_1 + Y_2 + Z_3) + \varepsilon^2(\dots) + \varepsilon^3(\dots)).$$

Wenn nun nicht $X_1 + Y_2 + Z_3 \geq 0$ eine notwendige Folge der obigen sechs Ungleichungen ist, so würden wir ε als positive Größe weiter so klein wählen können, daß hierbei Δ sich verringert. Für ein Minimum von Δ ist danach notwendig, daß die linke Seite dieser letzten Ungleichung eine homogene lineare Kombination der linken Seiten jener früheren sechs Ungleichungen mit nicht negativen Koeffizienten ist, d. h. bei einem Minimum von Δ müssen die drei Gleichungen

$$\begin{aligned}
 (23) \quad X &= lL - sS + tT, \\
 Y &= rR + mM - tT, \\
 Z &= -rR + sS + nN
 \end{aligned}$$

mit gewissen (und zwar durchweg nicht negativen) Faktoren l, m, n, r, s, t zu erfüllen sein.

32. Wir betrachten jetzt die Umstände des Falles (II) in 21.:

Es sollen die Punkte

$$\begin{array}{cccccc}
 \mathfrak{L}, & \mathfrak{M}, & \mathfrak{N}, & \mathfrak{R}, & \mathfrak{S}, & \mathfrak{T} \\
 1, 0, 0; & 0, 1, 0; & 0, 0, 1; & 0, 1, 1; & 1, 0, 1; & 1, 1, 0
 \end{array}$$

auf der Fläche $\mathfrak{F}$ und der Punkt $1, 1, 1$ außerhalb $\mathfrak{R}$ liegen.

Wir legen durch jeden der Punkte $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}, \mathfrak{R}, \mathfrak{S}, \mathfrak{T}$ eine Stützebene an $\mathfrak{R}$; die Gleichungen dieser Ebenen seien:

$$\begin{aligned}
 (24) \quad L &\equiv X - \lambda_2 Y - \lambda_3 Z = 1, \\
 M &\equiv -\mu_1 X + Y - \mu_3 Z = 1, \\
 N &\equiv -\nu_1 X - \nu_2 Y + Z = 1, \\
 R &\equiv -\varrho_1 X + \varrho_2 Y + (1 - \varrho_2)Z = 1, \\
 S &\equiv (1 - \sigma_3)X - \sigma_2 Y + \sigma_3 Z = 1, \\
 T &\equiv \tau_1 X + (1 - \tau_1)Y - \tau_3 Z = 1.
 \end{aligned}$$

In $\mathfrak{R}$ müssen die Ausdrücke L, M, N, R, S, T durchweg ≥ -1 und ≤ 1 bleiben. Die Form L wird für jene sechs Gitterpunkte bzw.

$$1, -\lambda_2, -\lambda_3, -\lambda_2 - \lambda_3, 1 - \lambda_3, 1 - \lambda_2;$$

mithin ist $0 \leq \lambda_2 \leq 1$, $0 \leq \lambda_3 \leq 1$ und $\lambda_2 + \lambda_3 \leq 1$. Die Form R wird für die sechs Punkte bzw.

$$-\varrho_1, \varrho_2, 1 - \varrho_2, 1, -\varrho_1 + 1 - \varrho_2, -\varrho_1 + \varrho_2;$$

mithin ist $0 \leq \varrho_2 \leq 1$ und entweder $0 \leq \varrho_1 \leq 1$ oder $0 \leq -\varrho_1$ und dabei zugleich $-\varrho_1 \leq \varrho_2$ und $-\varrho_1 \leq 1 - \varrho_2$. Entsprechende Umstände gelten für $\mu_3, \mu_1; \nu_1, \nu_2; \sigma_3, \sigma_2; \tau_1, \tau_3$.

Wir variieren nun $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}$ derart, daß ihre neuen Orte in bezug auf das alte X, Y, Z -Koordinatensystem werden:

$$1 + \varepsilon X_1, \varepsilon Y_1, \varepsilon Z_1; \quad \varepsilon X_2, 1 + \varepsilon Y_2, \varepsilon Z_2; \quad \varepsilon X_3, \varepsilon Y_3, 1 + \varepsilon Z_3,$$

wobei ε ein Parameter sei. Dabei soll jeder der sechs Punkte $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}, \mathfrak{R}, \mathfrak{S}, \mathfrak{T}$ in der für ihn konstruierten Stützebene an $\mathfrak{R}$ bleiben, d. h. es sollen die Gleichungen

$$L_1 = 0, M_2 = 0, N_3 = 0, R_2 + R_3 = 0, S_1 + S_3 = 0, T_1 + T_2 = 0$$

für die bezüglichen Systeme X_i, Y_i, Z_i gelten.

Falls nun auf der Fläche $\mathfrak{F}$ außer $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}, \mathfrak{R}, \mathfrak{S}, \mathfrak{T}$ und den Gegenpunkten keine weiteren Gitterpunkte vorhanden sind, treten bei hinreichend

kleinem Werte von $|\varepsilon|$ während der hier vorzunehmenden Variation des Gitters keine Gitterpunkte in $\mathfrak{R}$ ein, und zeigt eine ähnliche Überlegung wie in 31., daß für ein Minimum von Δ das Bestehen der drei Gleichungen

$$(25) \quad \begin{aligned} X &= lL + sS + tT, \\ Y &= rR + mM + tT, \\ Z &= rR + sS + nN \end{aligned}$$

mit irgendwelchen Faktoren l, m, n, r, s, t erforderlich ist.

Wofern jedoch noch weitere Gitterpunkte auf $\mathfrak{F}$ vorhanden sind, so können wir, wie nun gezeigt werden soll, die Formen $L, M, \dots, T$ jedenfalls immer derart wählen, daß während der fraglichen Variation bei hinreichend kleinem $|\varepsilon|$ kein Gitterpunkt ins Innere von $\mathfrak{R}$ eintritt, und finden wir dann wieder für ein Minimum von Δ die Gleichungen (25) als notwendig.

In der Tat, wir berücksichtigen die Bemerkung in 15. und gebrauchen zudem die folgenden Ungleichungen sowie die entsprechenden Beziehungen bei Permutation der Variablen:

$$F\left(-\frac{3}{1}, -\frac{2}{2}, \pm 4\right) + F(1, 1, 0) + F(0, 1, 0) \geq 4F\left(\frac{1}{0}, \frac{1}{0}, \pm 1\right),$$

$$F\left(-\frac{2}{1}, -\frac{2}{1}, \pm 3\right) + F(1, 1, 0) \geq 3F\left(\frac{1}{0}, \frac{1}{0}, \pm 1\right),$$

$$F(-a, -b, c) + aF(1, 0, 1) + bF(0, 1, 1) \geq (a + b + c)F(0, 0, 1),$$

$$(-a, -b, c = -1, -2, 3; -1, -2, 2; -1, -1, 2),$$

$$F(-2, 2, 3) + 2F(-1, -1, 0) + 3F(-1, 0, -1) \geq 7F(-1, 0, 0),$$

$$F(1, 2, 3) + F(1, 1, 0) + F(1, 0, 0) \geq 3F(1, 1, 1),$$

$$F(1, -1, 3) + F(0, 1, 1) + F(-1, 0, 0) \geq 4F(0, 0, 1),$$

$$F(1, 2, 2) + F(1, 0, 0) \geq 2F(1, 1, 1),$$

$$F(-1, 2, 2) + \frac{1}{2}F(1, 1, 0) + \frac{1}{2}F(1, 0, 1) \geq \frac{5}{2}F(0, 1, 1),$$

$$F(1, 1, 2) + F(1, 1, 0) \geq 2F(1, 1, 1),$$

$$F(0, -1, 2) + F(0, 1, 1) \geq 3F(0, 0, 1).$$

Danach können nur noch diejenigen Gitterpunkte als auf $\mathfrak{F}$ gelegen in Frage kommen, deren Koordinaten, von der Reihenfolge abgesehen, eines der folgenden Systeme ergeben

$$-1, 1, 2; 0, 1, 2; -1, 1, 1; 0, -1, 1.$$

Wenn $F(-1, 1, 2) = 1$ ist, so haben wir nach dem Ausdrucke von L in (24) notwendig $L = X$, und wir können daher $S = L$ und $T = L$ wählen; dann folgt oben $X_1 = 0, X_2 = 0, X_3 = 0$. Der Körper $\mathfrak{R}$ befindet sich ganz im Bereiche $-1 \leq X \leq 1$, und die in den begrenzenden

Ebenen $X = \pm 1$ gelegenen Gitterpunkte bleiben bei der vorzunehmenden Variation in diesen Ebenen. — Gleichzeitig könnte in der Ebene $X = 0$ als ein weiterer Gitterpunkt auf $\mathfrak{F}$ der Punkt $0, 1, 2$ oder $0, 2, 1$ oder $0, -1, 1$ auftreten; zu dem Ende müßte $\varrho_2 = 1$ bzw. $\varrho_2 = 0$ bzw. $\nu_2 = 0$ sein, und könnten wir dann $M = R$ bzw. $N = R$ bzw. $R = N$ wählen mit dem Erfolge, daß die anzuschließende Variation jenen Gitterpunkt nicht ins Innere von $\mathfrak{R}$ führt.

Wenn $F(-1, 1, 1) = 1$ ist, so finden wir wieder $L = X$ und treffen genau die nämlichen Bemerkungen wie soeben zu.

Wenn $F(0, 1, 2) = 1$ ist, so folgt mit Notwendigkeit $N = -Y + Z$, und können wir $S = N$ und $T = -N$ wählen, um den Zweck, den wir im Auge haben, zu erreichen. — Gleichzeitig könnte in der Ebene $-Y + Z = 0$ der Gitterpunkt $-1, 1, 1$ auf $\mathfrak{F}$ auftreten, in welchem Falle wir unter Vorwegnahme dieser Beziehung wie hier zuletzt verfahren.

Wenn $F(0, -1, 1) = 1$ ist, so folgt $\mu_3 = 0$, $\nu_2 = 0$ und bildet der Schnitt von $\mathfrak{R}$ mit der Ebene $X = 0$ das Viereck mit den Ecken $0, \pm 1, \pm 1$; wir können dann $R = N$ wählen, wobei mit $0, 0, 1$ und $0, 1, 1$ auch $0, -1, 1$ in der Ebene $N = 1$ verbleiben wird. — Finden sich zwei Punkte wie $1, -1, 0$ und $-1, 0, 1$ gleichzeitig auf $\mathfrak{F}$, so folgt $L = X$, und können wir wie vorhin im Falle $F(-1, 1, 1) = 1$ vorgehen.

33. Wir betrachten endlich den Fall (III) aus 21.: *Es sollen die sieben Gitterpunkte*

$$\mathfrak{L}, \quad \mathfrak{M}, \quad \mathfrak{N}, \quad \mathfrak{R}, \quad \mathfrak{S}, \quad \mathfrak{T}, \quad \mathfrak{Q}$$

$$1, 0, 0; 0, 1, 0; 0, 0, 1; 0, 1, 1; 1, 0, 1; 1, 1, 0 \text{ und } 1, 1, 1$$

auf $\mathfrak{F}$ liegen.

Wir gebrauchen in bezug auf die ersten 6 Punkte dieselben Bezeichnungen wie in 32., und ferner sei

$$(26) \quad Q \equiv \kappa_1 X + \kappa_2 Y + \kappa_3 Z = 1 \quad (\kappa_1 + \kappa_2 + \kappa_3 = 1)$$

die Gleichung einer Stützebene an $\mathfrak{R}$ durch den Punkt $\mathfrak{Q}$. Der Ausdruck Q wird für jene 7 Gitterpunkte bzw.

$$\kappa_1, \kappa_2, \kappa_3, \kappa_2 + \kappa_3, \kappa_1 + \kappa_3, \kappa_1 + \kappa_2, \kappa_1 + \kappa_2 + \kappa_3 = 1,$$

so daß $\kappa_1, \kappa_2, \kappa_3$ sämtlich ≥ 0 und ≤ 1 sind. Ferner wird für den Punkt $\mathfrak{Q}$ der Ausdruck $R = 1 - \varrho_1$, so daß jetzt notwendig $0 \leq \varrho_1 \leq 1$ ist, und analog für σ_2, τ_3 .

Wir variieren nun die Punkte $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}$ in

$$1 + \varepsilon X_1, \varepsilon Y_1, \varepsilon Z_1; \varepsilon X_2, 1 + \varepsilon Y_2, \varepsilon Z_2; \varepsilon X_3, \varepsilon Y_3, 1 + \varepsilon Z_3,$$

und es sollen dabei $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}, \mathfrak{Q}, \mathfrak{R}, \mathfrak{S}, \mathfrak{T}$ in den bezüglichen Stützebenen an $\mathfrak{R}$ verbleiben, also für die Systeme X_i, Y_i, Z_i die Relationen gelten

$$(27) \quad \begin{aligned} L_1 = 0, \quad M_2 = 0, \quad N_3 = 0, \quad Q_1 + Q_2 + Q_3 = 0, \\ R_2 + R_3 = 0, \quad S_1 + S_3 = 0, \quad T_1 + T_2 = 0. \end{aligned}$$

Eine ähnliche Überlegung wie in 32. zeigt, wofern bei hinreichend kleinem Betrage von ε durch jene Variationen kein Gitterpunkt ins Innere von $\mathfrak{R}$ eindringt, daß für ein Minimum von Δ das Bestehen der drei Gleichungen

$$(28) \quad \begin{aligned} X &= lL + sS + tT + qQ, \\ Y &= rR + mM + tT + qQ, \\ Z &= rR + sS + nN + qQ \end{aligned}$$

mit irgendwelchen Konstanten l, m, n, q, r, s, t notwendig ist. Die ausgesprochene Bedingung ist ohne weiteres erfüllt, falls $\mathfrak{F}$ überhaupt keine Gitterpunkte neben $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}, \mathfrak{Q}, \mathfrak{R}, \mathfrak{S}, \mathfrak{T}$ und den Gegenpunkten enthält, und anderenfalls kann ihr wenigstens immer durch geeignete Wahl der Ausdrücke L, M, N, Q, R, S, T genügt werden.

In der Tat, wie die in 32. entwickelten Ungleichungen dartun, könnten jetzt als Gitterpunkte auf $\mathfrak{F}$ außer den bereits dort betrachteten Punkten noch die Punkte 1, 2, 3; 1, 2, 2; 1, 1, 2 und die daraus durch Permutation der Koordinaten entstehenden Systeme in Betracht kommen. Vermöge der Substitution

$$X^* = -X, \quad Y^* = Y - X, \quad Z^* = Z - X$$

nun gehen die Wertsysteme X, Y, Z :

$$\begin{aligned} &1, 0, 0; 0, 1, 0; 0, 0, 1; 0, 1, 1; 1, 0, 1; 1, 1, 0; 1, 1, 1; \\ &1, 2, 3; 1, 2, 2; 1, 1, 2; -1, 1, 0; 0, -1, 1 \end{aligned}$$

in die Wertsysteme X^*, Y^*, Z^* :

$$\begin{aligned} &-1, -1, -1; 0, 1, 0; 0, 0, 1; 0, 1, 1; -1, -1, 0; -1, 0, -1; -1, 0, 0; \\ &-1, 1, 2; -1, 1, 1; -1, 0, 1; 1, 2, 1; 0, -1, 1 \end{aligned}$$

über, und es erledigt sich das Auftreten irgendeines der zuletzt genannten Gitterpunkte auf $\mathfrak{F}$ wesentlich wie in 32. das Auftreten eines der Punkte $-1, 1, 2; -1, 1, 1; -1, 0, 1$.

Als einzige neue Möglichkeit kommt in Frage, daß zwei Punkte wie 1, $-1, 0$ und 1, 1, 2 gleichzeitig sich auf $\mathfrak{F}$ vorfinden. Dazu müssen wir $\lambda_2 = 0, \mu_1 = 0, \nu_1 + \nu_2 = 1, \kappa_3 = 0$ haben. Wir können jedenfalls eine Relation

$$l_0 L + m_0 M + n_0 N = q_0 Q$$

mit Koeffizienten l_0, m_0, n_0, q_0 , die nicht sämtlich Null sind, herstellen. Da die Rollen der Paare $\mathfrak{L}, \mathfrak{M}$ und $\mathfrak{N}, \mathfrak{Q}'$ sowie der Elemente in jedem Paare hier vertauschbar sind, dürfen wir annehmen, es seien $|l_0|, |m_0|, |n_0|$ sämtlich $\leq |q_0|$. Durch Verwendung der Punkte 1, 1, 1 und 0, 0, -1 erhalten wir aus der letzten Gleichung

$$l_0(1 - \lambda_3) + m_0(1 - \mu_3) = q_0, \quad l_0\lambda_3 + m_0\mu_3 = n_0.$$

Für $\lambda_3 = 0$ würde $L = X$ folgen, welchen Fall wir wie in 32. erledigen könnten. Wir denken uns daher $\lambda_3 > 0$ und ebenso $\mu_3 > 0$. Als dann zeigen die zwei Gleichungen hier, daß l_0, m_0 und weiter n_0 dasselbe Vorzeichen wie q_0 haben, und wir richten $l_0 + m_0 = n_0 + q_0 = 1$ ein. Nun können wir

$$(29) \quad T = l_0 L + m_0 M = -n_0 N + q_0 Q$$

wählen. Dabei wird in der Tat $T = 1$ eine Stützebene an $\mathfrak{R}$ durch den Punkt $\mathfrak{X}$, und bei der in Rede stehenden Variation folgt durch die Gleichungen (27):

$$l_0 L_2 + m_0 M_1 = 0, \quad -n_0(N_1 + N_2 + N_3) - q_0 Q_3 = 0,$$

so daß einerseits nicht L_2 und M_1 , andererseits nicht $N_1 + N_2 + N_3$ und Q_3 zwei von Null verschiedene Werte mit gleichem Vorzeichen sein können. Nun wird für den Punkt 1, $-1, 0$ bei jener Variation: $L = 1 - \varepsilon L_2$, $-M = 1 - \varepsilon M_1$, so daß wenigstens eine dieser zwei Größen ≥ 1 bleibt, und für den varierten Punkt 1, 1, 2 wird $N = 1 + \varepsilon(N_1 + N_2 + N_3)$, $Q = 1 + \varepsilon Q_3$, so daß wenigstens eine dieser zwei Größen ≥ 1 bleibt.

§ 9. Dichteste Lagerung von Tetraedern.

34. Wir wenden jetzt unsere Ergebnisse speziell auf die dichteste Lagerung von *Tetraedern* oder von *Oktaedern* an. Es sei

$\varphi = -\xi + \eta + \zeta$, $\chi = \xi - \eta + \zeta$, $\psi = \xi + \eta - \zeta$, $\omega = -\xi - \eta - \zeta$,
so gilt

$$(30) \quad \varphi + \chi + \psi + \omega = 0$$

und definieren die Ungleichungen

$$\varphi \leq \frac{1}{4}, \quad \chi \leq \frac{1}{4}, \quad \psi \leq \frac{1}{4}, \quad \omega \leq \frac{1}{4}$$

den Bereich eines Tetraeders; in diesem Bereiche sind andererseits $\varphi, \chi, \psi, \omega$ stets $\geq -\frac{3}{4}$. Dieses Tetraeder nehmen wir jetzt für den Grundkörper K , sein Volumen ist $J = \frac{1}{24}$. Der zugehörige Körper $\frac{1}{2}(K + K')$ mit Mittelpunkt, der in bezug auf genau dieselben Gitter ($\mathfrak{G}$) wie K solche Anordnungen gestattet, wobei die Körper um die verschiedenen Gitterpunkte gesondert liegen, ist dann das Oktaeder

$$-\frac{1}{2} \leq \varphi \leq \frac{1}{2}, \quad -\frac{1}{2} \leq \chi \leq \frac{1}{2}, \quad -\frac{1}{2} \leq \psi \leq \frac{1}{2}, \quad -\frac{1}{2} \leq \omega \leq \frac{1}{2};$$

diese acht Ungleichungen können in die eine Bedingung

$$|\xi| + |\eta| + |\zeta| \leq \frac{1}{2}$$

zusammengefaßt werden. Das Volumen dieses Oktaeders $\frac{1}{2}(K + K')$ ist $\frac{1}{6}$. Wir substituieren nun

$\xi = \alpha_1 X + \alpha_2 Y + \alpha_3 Z$, $\eta = \beta_1 X + \beta_2 Y + \beta_3 Z$, $\zeta = \gamma_1 X + \gamma_2 Y + \gamma_3 Z$, wobei die Determinante der drei linearen Ausdrücke $\neq 0$ sei, und wir fragen nach dem Minimum für den absoluten Betrag Δ dieser Determinante, während vom ganzen Zahlengitter in X, Y, Z bloß der Nullpunkt ins Innere des Oktaeders $\mathfrak{R} = K + K'$ fällt.

Wir bringen nacheinander die verschiedenen Regeln des § 8 zur Anwendung. Von dem einen in 33. am Schlusse berührten Ausnahmefalle abgesehen, der, wie sich zeigen wird, hier nicht in Betracht kommt, können wir die jedesmal einzuführenden 6 oder 7 Stützebenen an $\mathfrak{R}$ immer als Seitenflächen dieses Oktaeders gewählt voraussetzen, wir können also annehmen, daß jede einzelne der Formen $L, M, N, \dots$ mit einem von den Ausdrücken $\pm \varphi, \pm \chi, \pm \psi, \pm \omega$ übereinstimmt.

35. Wir wollen vorweg einen speziellen Fall behandeln, wie er in 32. zur Sprache kam. Es mögen die Umstände des Falles (II) dort gelten, und dabei sei $\omega = Z$. Der Schnitt von $\mathfrak{R}$ mit der Ebene $Z = 0$ ist alsdann ein Sechseck mit O als Mittelpunkt, bei dem die Diagonalen den Seiten parallel sind. Auf diesem Sechseck liegen die Gitterpunkte $\mathfrak{L}, \mathfrak{M}, \mathfrak{T}(1, 0, 0; 0, 1, 0; 1, 1, 0)$. Wir können annehmen, daß die Ausdrücke L, M, T , in irgendeiner Reihenfolge genommen, mit $\pm \varphi, \pm \chi, \pm \psi$ übereinstimmen, denn wenn zwei jener Punkte in einer und derselben Seite des Sechsecks liegen sollten, so müßten die drei Punkte sämtlich Ecken des Sechsecks werden.

Die Gleichung (30) ergibt nun mit Rücksicht auf die Ausdrücke (24): $\mu_1 = 1 - \tau_1$, $\lambda_2 = \tau_1$, und die Gleichungen (25) führen dann zu $\tau_1 = \frac{1}{2}$, so daß die 3 Gitterpunkte $\mathfrak{L}, \mathfrak{M}, \mathfrak{T}$ mit ihren Gegenpunkten die Mitten der Seiten des Sechsecks bilden. Durch kontinuierliche Variation von $\mathfrak{N}(0, 0, 1)$ auf der Seitenfläche $Z = 1$, wobei Δ sich nicht ändert, können wir insbesondere folgende Ausdrücke erzielen:

$$L = X - \frac{1}{2}Y - \frac{1}{2}Z, \quad M = -\frac{1}{2}X + Y - \frac{1}{2}Z, \quad T = \frac{1}{2}X + \frac{1}{2}Y,$$

$$N = R = S = Z;$$

dabei liegen in der Seitenfläche $Z = 1$ von $\mathfrak{R}$ die fünf Gitterpunkte

$$-1, -1, 1; 0, 0, 1; 1, 1, 1; 1, 0, 1; 0, 1, 1.$$

Variieren wir nun $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}$ in

$$1 + \varepsilon, \varepsilon, \varepsilon; -\varepsilon, 1 - \varepsilon, -\varepsilon; -\varepsilon, -2\varepsilon, 1,$$

während ε positiv sei, so tritt bei hinreichend kleinem ε kein Gitter-

punkt ins Innere von $\mathfrak{R}$ ein und Δ geht in $\Delta(1 - \varepsilon^2)$ über. Also liegt hier kein Minimum von Δ vor.

36. Nach Beseitigung dieses speziellen Falles gehen wir zuerst auf die Umstände des Falles (I) gemäß 31. ein, wobei die Punkte $1, 0, 0$; $0, 1, 0$; $0, 0, 1$; $0, 1, -1$; $-1, 0, 1$; $1, -1, 0$ auf $\mathfrak{F}$ und $-1, 1, 1$; $1, -1, 1$; $1, 1, -1$ außerhalb $\mathfrak{R}$ liegen sollen.

Da für die 6 Ausdrücke L, M, N, R, S, T in (22) nur die vier Paare $\pm \varphi, \pm \chi, \pm \psi, \pm \omega$ in Betracht kommen, so müssen unter diesen Ausdrücken entweder irgend drei oder zweimal je zwei anzugeben sein, die bis auf den Faktor ± 1 übereinstimmen.

Nach den Bedingungen für die Koeffizienten in (22) kann nicht $L = -M$ sein. — Die Annahme $L = R$ führt zu $L = X + Y$. Der Schnitt von $\mathfrak{R}$ mit der Ebene $X + Y = 0$ ist dann ein Sechseck mit nur zwei Paaren von Gitterpunkten auf dem Rande, durch $\mathfrak{N}$ und $\mathfrak{T}$ repräsentiert, und wir können, sei es durch alleinige Variation von $\mathfrak{N}$ oder von $\mathfrak{T}$, den Wert von Δ verringern.

Stellen wir uns die Figur des Tetraeders aus den Vektoren $1, 0, 0$; $0, 1, 0$; $0, 0, 1$; $0, 1, -1$; $-1, 0, 1$; $1, -1, 0$ vor (Fig. 3 links in 19.) und beachten, wie darin die Rollen der einzelnen Vektoren vertauscht werden können, so leuchtet ein, daß wir nach den letzten Bemerkungen jetzt überhaupt die Gleichheit für irgend zwei der Ausdrücke $\pm L, \dots, \pm T$ ausschließen können, falls die zugehörigen Systeme $\mathfrak{Q}, \mathfrak{Q}', \dots, \mathfrak{T}, \mathfrak{T}'$ zwei solchen Vektoren in jenem Tetraeder entsprechen, die entweder sich in einer Ecke mit ihren Richtungen aneinanderschließen oder auf gegenüberliegenden Kanten Platz finden. Danach brauchen wir nur noch die folgenden Annahmen zu diskutieren:

1) $L = M = N$. — Daraus folgt $L = X + Y + Z$. Es sei dieser Ausdruck etwa $= \omega$, so ergeben sich ganz entsprechende Umstände wie in dem Falle $\omega = Z$, der in 35. vorweggenommen wurde, und ist Δ hier nicht ein Minimum.

2) $L = M, R = -S$. — Hierbei folgt

$$L = X + Y + \lambda_3 Z, \quad R = \varrho_2(X + Y) - (1 - \varrho_2)Z.$$

Aus der dritten Gleichung in (23):

$$Z = -rR + sS + nN$$

geht dann entweder $R = -Z$ hervor, welcher Fall sich ähnlich wie in 35. erledigt, oder $N \equiv 0 \pmod{X + Y, Z}$. In letzterem Falle wäre die Determinante von L, R, N Null, während wir uns diese Formen hier als drei verschiedene Ausdrücke $\pm \varphi, \pm \chi, \pm \psi, \pm \omega$ denken müssen, um nicht auf schon erledigte Fälle zurückzukommen.

3) $L = N$, $R = -S$. — Hier wird

$$N = X + \nu_2 Y + Z, \quad R = \varrho_2 (X + Y) - (1 - \varrho_2) Z,$$

und die soeben erwähnte Gleichung für Z erfordert entweder $N = X + Y + Z$, oder $R = -Z$, welche Fälle schon Erledigung fanden.

37. Wir betrachten zweitens die Umstände des Falles (II) gemäß 32., wobei die Punkte 1, 0, 0; 0, 1, 0; 0, 0, 1; 0, 1, 1; 1, 0, 1; 1, 1, 0 auf $\mathfrak{F}$ und 1, 1, 1 außerhalb $\mathfrak{R}$ liegen sollen.

Nach den Bedingungen für die Koeffizienten in den Ausdrücken (24) kann nicht $L = M$ sein. — Die Annahme $L = -M$ führt zu $L = X - Y$ und ist hier wesentlich wie der in 35. vorausgeschickte Spezialfall zu erledigen. — Ferner kann nicht $L = R$, nicht $L = -S$ und nicht $L = -T$ sein. — Die Annahme $R = -S$ führt auf $R = -X + Y$.

Mit Rücksicht auf die Vertauschbarkeit der Rollen der drei Paare L, R ; M, S ; N, T bleiben hiernach nur die folgenden verschiedenen Fälle zu diskutieren, wobei wir bei der Behandlung jedes einzelnen Falles annehmen, daß die Umstände der vorhergegangenen Fälle ausgeschlossen sind:

1) $R = S = N$. — Hier folgt $N = Z$ und wäre dieses der Fall aus 35.

2) $R = S = T$. — Wir haben danach $R = \frac{1}{2} (X + Y + Z)$. Die Relationen (25) führen zu $\lambda_2 = \lambda_3$, $\mu_1 = \mu_3$, $\nu_1 = \nu_2$. Gemäß (30) muß dann $\pm L \pm M \pm N = R$ sein, und diese Bedingung zieht notwendig die Ausdrücke

$$L = X - \frac{1}{4} Y - \frac{1}{4} Z, \quad M = -\frac{1}{4} X + Y - \frac{1}{4} Z, \quad N = -\frac{1}{4} X - \frac{1}{4} Y + Z$$

nach sich. Variieren wir dann die Punkte $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}$ in

$$1, \varepsilon, -\varepsilon; \varepsilon, 1, -\varepsilon; 0, 0, 1,$$

so tritt bei hinreichend kleinem Betrage von ε kein Gitterpunkt ins Innere von $\mathfrak{R}$ ein, und Δ geht in $\Delta(1 - \varepsilon^2)$ über; also ist hier kein Minimum von Δ vorhanden.

3) $R = -L$, $S = -M$. — Wir erhalten

$$L = X - \lambda_2 Y - (1 - \lambda_2) Z, \quad M = -\mu_1 X + Y - (1 - \mu_1) Z.$$

Für den Punkt 1, 1, 1 ist $L = 0$, $M = 0$ und muß für ihn daher $N \neq 0$ sein, da wir in L, M, N drei der Ausdrücke $\pm \varphi, \pm \chi, \pm \psi, \pm \omega$ haben; mithin folgt $\nu_1 + \nu_2 < 1$ und fällt daher N auch von $-T$ verschieden aus. Wir müssen nun $\pm L \pm M \pm N = T$ haben. Da 1, 1, 1 nicht im Inneren von $\mathfrak{R}$ liegt, folgt daraus notwendig $\nu_1 = 0$, $\nu_2 = 0$, also $N = Z$.

4) $R = -L$, $S = N$. — Hier folgt

$$-R = X - \lambda_2 Y - (1 - \lambda_2) Z, \quad N = -\nu_2 Y + Z,$$

und die dritte Gleichung in (25):

$$Z = rR + sS + nN$$

erfordert $n = Z$.

5) $R = M$, $S = L$. — Wir haben dann

$$M = -\mu_1 X + Y, \quad L = X - \lambda_2 Y,$$

und aus der ersten Gleichung in (25): $X = lL + sS + tT$ würde $L = X$ oder $T = \tau_1 X + (1 - \tau_1) Y$ folgen. In letzterem Falle aber wäre die Determinante von L, M, T Null, während diese Formen hier drei der Ausdrücke $\pm \varphi, \pm \chi, \pm \psi, \pm \omega$ darstellen sollen.

6) $R = M$, $S = N$. — Hier wäre

$$R = -\mu_1 X + Y, \quad N = -\nu_2 Y + Z,$$

und die im Falle 4) genannte Gleichung für Z würde $N = Z$ oder $R = Y$ erfordern.

7) $R = S$, $T = N$. — Wir erhalten hiernach

$$N = -\nu_1 X - (1 - \nu_1) Y + Z, \quad R = \varrho_2 (X + Y) + (1 - \varrho_2) Z, \quad \varrho_2 \leq \frac{1}{2},$$

und aus der Relation für Z folgt dann $R = Z$ oder aber $N = -\frac{1}{2}X - \frac{1}{2}Y + Z$.

In letzterem Falle führen die Gleichungen (25) und die Beziehung $\pm L \pm M \pm N = R$ weiter zu den Ausdrücken

$$L = X - \frac{1}{4}Y - \frac{1}{8}Z, \quad M = -\frac{1}{4}X + Y - \frac{1}{8}Z,$$

$$R = S = \frac{1}{4}X + \frac{1}{4}Y + \frac{3}{4}Z.$$

Variieren wir nun die Punkte $\mathfrak{L}, \mathfrak{M}, \mathfrak{N}$ in

$$1 - \frac{1}{2}\varepsilon, -\frac{3}{2}\varepsilon, -\varepsilon; -\frac{3}{2}\varepsilon, 1 - \frac{1}{2}\varepsilon, -\varepsilon; \varepsilon, \varepsilon, 1 + \varepsilon,$$

so tritt bei hinreichend kleinem Betrage von ε kein Gitterpunkt ins Innere von $\mathfrak{K}$ ein, und Δ geht in $\Delta(1 - \varepsilon^2)$ über. Also ist hier Δ nicht ein Minimum.

8) $R = S$, $T = L$. — Hier wird

$$S = \varrho_2 (X + Y) + (1 - \varrho_2) Z, \quad L = X - \lambda_3 Z,$$

und aus $X = lL + sS + tT$ folgt notwendig $L = X$ oder $S = Z$.

Auf diese Weise hat auch die Betrachtung des Falles (II) zu keinem Falle eines Minimums von Δ geführt.

38. Wir betrachten endlich die Umstände des Falles (III) gemäß 33., wobei auf $\mathfrak{F}$ die sieben Punkte $1, 0, 0$; $0, 1, 0$; $0, 0, 1$; $0, 1, 1$; $1, 0, 1$; $1, 1, 0$; $1, 1, 1$ liegen sollen.

Die Annahme $R = S$ würde jetzt, da hier ϱ_1 und $\sigma_2 \geq 0$ sind, $R = Z$ zur Folge haben. — Beachten wir noch, daß die Rollen der vier Punkte

$\mathfrak{L}, \mathfrak{M}, \mathfrak{N}, \mathfrak{D}'$ durchaus vertauschbar sind, so können wir uns auf die Diskussion folgender Annahmen beschränken und dabei ferner $L, M, N, -Q$ als mit $\pm \varphi, \pm \chi, \pm \psi, \pm \omega$ identisch voraussetzen:

1) $L = -R, M = -S, N = -T$. — Hier würde für den Punkt 1, 1, 1 sich $L = 0$ herausstellen und ebenso $M = 0, N = 0$, was unmöglich ist.

2) $L = S, M = R, N = -T$. — Hier hätten wir

$$L = X - \lambda_2 Y, M = -\mu_1 X + Y, N = -\nu_1 X - (1 - \nu_1) Y + Z$$

und aus $\pm L \pm M \pm N = Q$ und nach (26) folgt notwendig $Q = Z$.

3) $L = S, M = T, N = R$. — In diesem Falle erhalten wir zunächst

$$L = X - \lambda_2 Y, M = Y - \mu_3 Z, N = -\nu_1 X + Z.$$

Würde 0, 1, -1 auf $\mathfrak{F}$ liegen, so müßte $N = Z$ sein; der am Schlusse von 33. behandelte Ausnahmefall kommt danach hier nicht in Frage.

Aus $\pm L \pm M \pm N = Q$ folgt $(1 - \nu_1) + (1 - \lambda_2) + (1 - \mu_3) = 1$. Die Relationen (28) in 33. ergeben sodann $\lambda_2 = \mu_3 = \nu_1$. Mithin kommen wir hier wesentlich zu folgenden Ausdrücken für die Formen $\varphi, \chi, \psi, \omega$:

$$(31) \quad \begin{aligned} \varphi &= X - \frac{2}{3} Y, \chi = Y - \frac{2}{3} Z, \psi = -\frac{2}{3} X + Z, \\ \omega &= -\frac{1}{3} X - \frac{1}{3} Y - \frac{1}{3} Z. \end{aligned}$$

Bei diesen Ausdrücken wird $\Delta = \frac{1}{4} \left(1 - \frac{8}{27}\right) = \frac{19}{108}$. Daß in diesem einzig übrig gebliebenen Falle Δ in der Tat ein Minimum sein muß, versteht sich von selbst und kann leicht verifiziert werden.

39. Beachten wir noch, daß in den Ausdrücken (31) für $\varphi, \chi, \psi, \omega$ die eine Form ω und für die anderen die zyklische Folge φ, χ, ψ bevorzugt erscheint, so können wir endlich das Resultat aussprechen:

Gegebene Tetraeder (Oktaeder) in unendlicher Anzahl, die sämtlich mit einem unter ihnen, dem Tetraeder

$$-\xi + \eta + \zeta \leq \frac{1}{4}, \xi - \eta + \zeta \leq \frac{1}{4}, \xi + \eta - \zeta \leq \frac{1}{4}, -\xi - \eta - \zeta \leq \frac{1}{4}$$

(dem Oktaeder

$$|\xi| + |\eta| + |\zeta| \leq \frac{1}{2})$$

kongruent und parallel orientiert sind, können sich auf acht Arten in dichtester gitterförmiger Lagerung befinden. Diese Lagerungen werden erhalten, indem $\pm \xi, \pm \eta, \pm \zeta$ oder $\pm \xi, \pm \zeta, \pm \eta$ irgendwie mit solchen Vorzeichen, deren Produkt +1 ist, gleich

$$-\frac{2}{6} X + \frac{3}{6} Y + \frac{1}{6} Z, \quad \frac{1}{6} X - \frac{2}{6} Y + \frac{3}{6} Z, \quad \frac{3}{6} X + \frac{1}{6} Y - \frac{2}{6} Z$$

gesetzt werden und das Zahlengitter in X, Y, Z als Ort für die Schwer-

punkte der Körper genommen wird. Der von den Tetraedern (Oktaedern) erfüllte Raum verhält sich dabei zu dem von ihnen freigelassenen Raume bzw. zu dem ganzen unendlichen Raume wie

$$\frac{1}{24} \left(\frac{1}{6} \right) : \frac{29}{216} \left(\frac{1}{108} \right) : \frac{19}{108}.$$

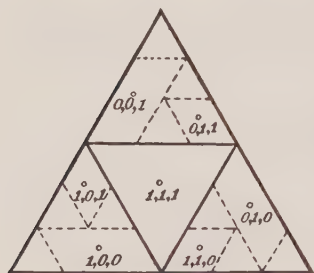


Fig. 6.

Die Fig. 6 zeigt für den durch die Ausdrücke (31) charakterisierten Fall der dichtesten Lagerung von Tetraedern oder Oktaedern die Lage der Gitterpunkte in dem halben, durch die Flächen $\varphi = 1$, $\chi = 1$, $\psi = 1$, $\omega = -1$ gebildeten Netze des Oktaeders $\mathfrak{R}$.

40. Wir geben diesem Satze ferner die folgende rein arithmetische Einkleidung:

Sind ξ, η, ζ drei lineare Formen in den Variablen x, y, z mit beliebigen reellen Koeffizienten und einer von Null verschiedenen Determinante, deren Betrag Δ ist, so kann man stets für x, y, z solche ganzzahlige Werte, die nicht sämtlich Null sind, finden, daß

$$|\xi| + |\eta| + |\zeta|,$$

also der Betrag jedes der vier Ausdrücke

$$-\xi + \eta + \zeta, \quad \xi - \eta + \zeta, \quad \xi + \eta - \zeta, \quad -\xi - \eta - \zeta$$

dabei $\leq \sqrt[3]{\frac{108}{19}} \Delta$ ausfällt.

An Stelle des Zeichens $\leq$ hier genügt das Zeichen $<$, falls ξ, η, ζ nicht gerade so beschaffen sind, daß $\pm \xi, \pm \eta, \pm \zeta$, mit irgendwelchen Vorzeichen und in irgendwelcher Reihenfolge, durch eine lineare ganzzahlige Substitution mit einer Determinante ± 1 in die Ausdrücke

$$-\frac{2}{6}X + \frac{3}{6}Y + \frac{1}{6}Z, \quad \frac{1}{6}X - \frac{2}{6}Y + \frac{3}{6}Z, \quad \frac{3}{6}X + \frac{1}{6}Y - \frac{2}{6}Z$$

transformiert werden können.

41. Es seien ξ, η, ζ wiederum drei lineare Formen in x, y, z mit beliebigen reellen Koeffizienten und einer von Null verschiedenen Determinante, deren Betrag Δ ist. Setzen wir

$$\omega_1 = \xi + \zeta, \quad \omega_2 = -\xi + \zeta, \quad \omega_3 = \eta - \zeta, \quad \omega_4 = -\eta - \zeta,$$

so entsteht

$$\omega_1 + \omega_2 + \omega_3 + \omega_4 = 0.$$

Eine Anwendung des Satzes aus 40. auf diese vier Ausdrücke $\omega_1, \omega_2, \omega_3, \omega_4$ ergibt, daß es stets möglich ist, für x, y, z ganzzahlige, von 0, 0, 0 verschiedene Werte zu finden, wofür

$$|\xi| + |\zeta|, \quad |\eta| + |\zeta|$$

beide $\leq \sqrt[3]{\frac{54}{19}} \Delta$ ausfallen.

Mit Hilfe der Ungleichungen

$$\frac{|\xi^2 \zeta|}{4} \leq \left(\frac{2 \left| \frac{\xi}{2} \right| + |\zeta|}{3} \right)^3, \quad \frac{|\eta^2 \zeta|}{4} \leq \left(\frac{2 \left| \frac{\eta}{2} \right| + |\zeta|}{3} \right)^3$$

folgt daraus weiter

$$|\xi^2 \zeta| \leq \frac{8}{19} \Delta, \quad |\eta^2 \zeta| \leq \frac{8}{19} \Delta.$$

Wir bemerken noch, daß in dem Grenzfalle, für den allein in dem letzten Satze das Zeichen = nötig war, das betreffende System x, y, z hier immer derart ausgesucht werden kann, daß dafür weder $\xi = \pm 2\zeta$ noch $\eta = \pm 2\zeta$ gilt, und wird dadurch in diesen weiteren Ungleichungen das Zeichen = entbehrlich.

Wenden wir das Ergebnis insbesondere auf die Ausdrücke

$$\xi = x - az, \quad \eta = y - bz, \quad \zeta = \frac{z}{t}$$

an, wobei a, b zwei beliebige reelle Größen sind und t ein positiver Parameter ($> \frac{54}{19}$) sei, so kommen wir zu dem Satze:

Sind a, b irgend zwei reelle Größen, so kann man stets solche ganze Zahlen x, y, z finden, daß $z > 0$, $|x - az|$, $|y - bz|$ beliebig klein und dabei

$$\left| \frac{x}{z} - a \right| < \sqrt{\frac{8}{19}} \cdot \frac{1}{z^{\frac{3}{2}}}, \quad \left| \frac{y}{z} - b \right| < \sqrt{\frac{8}{19}} \cdot \frac{1}{z^{\frac{3}{2}}}$$

sind.

Die Konstante $\sqrt{\frac{8}{19}}$ ist $= 0,648 \dots$

§ 10. Bemerkung zu einem Aufsätze von Lord Kelvin.

42. In den Baltimore lectures, appendix H, S. 618 ff. behandelt Lord Kelvin die dichtesten gitterförmigen Lagerungen von kongruenten und parallel orientierten konvexen Körpern und geht dabei von der Annahme aus, daß in einer dichtesten Lagerung vier sich gegenseitig berührende Körper auftreten. Dieser Umstand ereignet sich in der Tat im Falle (I) (vgl. 21., also z. B. für Kugeln) bei den Körpern, die sich um die vier Punkte 0, 0, 0; 1, 0, 0; 0, 1, 0; 0, 0, 1 lagern; im Falle (III) (also z. B. für Oktaeder) bei den Körpern um die vier Punkte 0, 0, 0; 1, 0, 0; 1, 1, 0; 1, 1, 1. Die Annahme ist aber nicht zutreffend im Falle (II).

Dieser Fall (II) tritt z. B. ein, wenn das von den 12 Ebenen

$$\begin{aligned}
X - \delta Y - \delta Z &= \pm \frac{1}{2}, \\
-\delta X + Y - \delta Z &= \pm \frac{1}{2}, \\
-\delta X - \delta Y + Z &= \pm \frac{1}{2}, \\
\varepsilon X + \frac{1}{2} Y + \frac{1}{2} Z &= \pm \frac{1}{2}, \\
\frac{1}{2} X + \varepsilon Y + \frac{1}{2} Z &= \pm \frac{1}{2}, \\
\frac{1}{2} X + \frac{1}{2} Y + \varepsilon Z &= \pm \frac{1}{2}
\end{aligned}$$

begrenzte Polyeder, wobei $0 < \delta < \frac{1}{2}$ und $0 < \varepsilon < \frac{1}{2}$ sei, den Grundkörper abgibt; bei hinreichend kleinen Werten von δ und von ε ist leicht zu sehen, daß dieses Polyeder durch die Parallelverschiebungen vom Nullpunkte nach allen Punkten mit ganzzahligen Koordinaten X, Y, Z ein System von Körpern in dichtester gitterförmiger Lagerung erzeugt.

XX.

Zur Geometrie der Zahlen.

(Mit Projektionsbildern auf einer Doppeltafel.)

(Verhandlungen des III. Internationalen Mathematiker-Kongresses. Heidelberg 1904.
S. 164—173.)

Im folgenden möchte ich versuchen, in kurzen Zügen einen Bericht über ein eigenartiges, zahlreicher Anwendungen fähiges Kapitel der Zahlentheorie zu geben, ein Kapitel, von dem Charles Hermite einmal als der „introduction des variables continues dans la théorie des nombres“ gesprochen hat. Einige hervorstechende Probleme darin betreffen die Abschätzung der kleinsten Beträge kontinuierlich veränderlicher Ausdrücke für ganzzahlige Werte der Variablen.

Die in dieses Gebiet fallenden Tatsachen sind zumeist einer geometrischen Darstellung fähig, und dieser Umstand ist für die in letzter Zeit hier erzielten Fortschritte derart maßgebend gewesen, daß ich geradezu das ganze Gebiet als die *Geometrie der Zahlen* bezeichnet habe.

(Fig. 1.) Die erste Figur illustriert für die Ebene dasjenige Theorem, welches mit Recht als das *Fundamentaltheorem* der Geometrie der Zahlen bezeichnet werden kann, weil es fast in jede Untersuchung auf diesem Gebiete hineinspielt.

In der Ebene denken wir uns irgendwelche Parallelkoordinaten x, y eingeführt, wobei noch für jede Achse der Einheitsmaßstab beliebig gewählt sein kann. Die Punkte mit ganzzahligen Koordinaten x, y bilden das *Zahlengitter* in x, y . Dieses Gitter kann auf mannigfaltige Weise durch ein Gerüst von kongruenten homothetischen Parallelogrammen gestützt werden, welche die Ebene lückenlos überdecken und als deren Ecken die Punkte des Gitters erscheinen. Bei einer vollständigen und einfachen Überdeckung der Ebene kommt danach sozusagen auf jeden Gitterpunkt ein homologes Gebiet von einem Flächeninhalt $\iint dx dy = 1$.

Nun denken wir uns irgendeine geschlossene konvexe Kurve, welche im Nullpunkt einen Mittelpunkt haben soll; sie darf auch geradlinige Stücke aufweisen. Das von ihr umschlossene Gebiet dilatieren wir vom

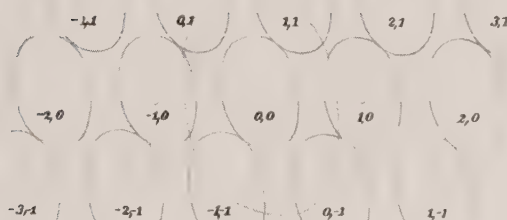
Nullpunkte aus (bzw. wir ziehen es zusammen) nach allen Richtungen in gleichem Verhältnisse zu einer solchen homothetischen, d. h. ähnlichen und ähnlich gelegenen, der grün umrandeten Figur, welche im Inneren den Nullpunkt als einzigen Gitterpunkt enthalten, ihren Rand aber durch wenigstens einen weiteren Gitterpunkt schicken soll. Ziehen wir weiter diese grüne Figur vom Nullpunkte aus im Verhältnisse $1:2$ zusammen und konstruieren um jeden anderen Gitterpunkt das homologe Gebiet, so erhalten wir offenbar lauter solche Gebiete, die nicht ineinander eindringen. Da nun bei lückenloser Überdeckung der Ebene im Durchschnitt auf einen Gitterpunkt ein Flächeninhalt $= 1$ kommt, so muß der Flächeninhalt eines solchen rot umgrenzten Gebiets < 1 bzw. $= 1$ sein, falls diese roten Gebiete ebenfalls die Ebene lückenlos ausfüllen. Das von der grünen Kurve umschlossene Gebiet hat daher einen Inhalt ≤ 4 . Treiben wir es endlich zu einem homothetischen Gebiete genau vom Inhalte 4 auf, so kommen wir zu dem Fundamentaltheoreme:

Ein konvexes Gebiet mit einem Mittelpunkt im Nullpunkt vom Flächeninhalt 4 enthält stets wenigstens einen weiteren Gitterpunkt.

Die Formeln in der Figur geben den analytischen Ausdruck dieses Theorems. Genügt eine Funktion $f(x, y)$ den Bedingungen (1)—(4) und ist J der Flächeninhalt des Gebiets $f \leq 1$, so kann man stets solche ganze Zahlen x, y , die nicht beide Null sind, finden, wofür die Funktion $f \leq 2J^{-\frac{1}{2}}$ ausfällt. (1) besagt, daß f wesentlich positiv, (2), daß es homogen von der ersten Dimension ist, (3), daß das Gebiet $f \leq 1$ einen Mittelpunkt hat, und wird die Länge der Verbindungsstrecke zweier Punkte mit den relativen Koordinaten x, y durch $f(x, y)$ definiert, so besagt (4), daß in einem Dreiecke die Summe zweier Seitenlängen niemals kleiner als die Länge der dritten Seite sein soll.

(Fig. 2.) Eine ausgezeichnete Anwendung dieses Satzes erläutert die nächste Figur. Bekanntlich tragen die gewöhnlichen Kettenbruchentwicklungen für Funktionen einer Variablen einen einfacheren Charakter als die analogen Entwicklungen für reelle Größen. Eine Funktion $f(z)$, die für $z = \infty$ einen Pol hat, läßt sich in einen Kettenbruch (1) entwickeln, worin $F_0(z)$ und die Nenner $F_1(z), F_2(z), \dots$ ganze rationale Funktionen sind. Die Näherungsbrüche dieses Kettenbruchs lassen sich von vornherein charakterisieren, ohne daß es nötig wäre, sie erst sukzessive durch die Entwicklung ausfindig zu machen. Als Näherungsbruch tritt hier jeder solche Quotient $P(z)/Q(z)$ zweier teilerfremden ganzen Funktionen auf, für welchen der Ausdruck (2), nach fallenden Potenzen von z entwickelt, mit einer Potenz von negativem Exponenten beginnt. Während eine sehr weitgehende Analogie zwischen den Eigenschaften der ganzen

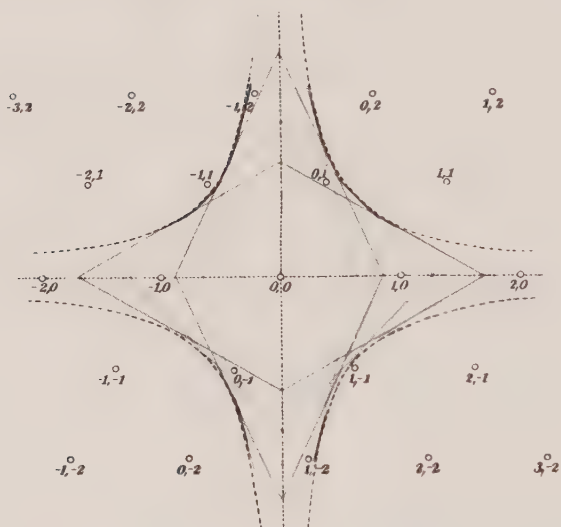
Fig. 1. Zahlengitter und konvexe Kurven.



$$f(x, y):$$

- (1) $f(x, y) > 0, x, y \neq 0, 0; f'(0, 0) = 0,$
- (2) $f(tx, ty) = tf(x, y), t > 0,$
- (3) $f(-x, -y) = f(x, y),$
- (4) $f(x_1, y_1) + f(x_2, y_2) \geq f(x_1 + x_2, y_1 + y_2),$
- (5) $f(x, y) \leq 1, \iint dx dy = J;$
- (6) $f(x, y) \leq \frac{2}{\sqrt{J}}.$

Fig. 2. Diagonalketten.



$$(1) \quad f(z) = c_m z^m + \dots + c_0 + \frac{c_1}{z} + \frac{c_2}{z^2} + \dots$$

$$= F_0(z) - \frac{1}{F_1(z)} - \frac{1}{F_2(z)} - \dots;$$

$$(2) \quad (P(z) - f(z)Q(z))Q(z).$$

$$(3) \quad \left| \frac{x}{y} - a \right| < \frac{1}{2y^2}, \quad (4) \quad a = g_0 - \frac{1}{g_1} - \frac{1}{g_2} - \dots$$

$$(5) \quad \xi = \alpha x + \beta y, \quad \eta = \gamma x + \delta y, \quad \alpha\delta - \beta\gamma = 1;$$

$$(6) \quad -\frac{1}{2} < \xi\eta < \frac{1}{2}.$$

Funktionen und der ganzen Zahlen besteht, schien zu dem genannten Satze ein entsprechender in der Arithmetik zu fehlen. Dieses Analogon finden wir darin, daß für eine beliebige reelle GröÙe a die sämtlichen gekürzten Brüche x/y , welche die Ungleichung (3) erfüllen, für welche also $(x - ay)y$ in den Grenzen $-\frac{1}{2}$ und $\frac{1}{2}$ liegt, sich als die Näherungsbrüche einer bestimmten Kettenbruchentwicklung (4) mit ganzzahligem g_0 und lauter positiven ganzzahligen $g_1, g_2, \dots$ anordnen.*)

Sind, um etwas allgemeinere Umstände zu betrachten, ξ, η zwei binäre lineare Formen in x, y mit beliebigen reellen Koeffizienten und einer Determinante $= 1$ (im speziellen würden wir $\xi = x - ay, \eta = y$ annehmen), so besteht ein sehr anschaulicher Zusammenhang zwischen den sämtlichen möglichen Auflösungen der Ungleichung $|\xi\eta| < \frac{1}{2}$ in ganzen Zahlen x, y ohne gemeinsamen Teiler.***) Wir zeichnen die Geraden $\xi = 0, \eta = 0$, etwa rechtwinklig zueinander, und die beiden Hyperbeln $\xi\eta = \frac{1}{2}$ und $\xi\eta = -\frac{1}{2}$. Wir legen eine beliebige Tangente an den Hyperbelast im ersten ξ, η -Quadranten und konstruieren dazu die Spiegelbilder in den drei anderen Quadranten, so daß wir ein Tangentenparallelogramm mit den Diagonalen $\xi = 0, \eta = 0$ erhalten. Ein solches Parallelogramm enthält nun stets wenigstens eine primitive (d. h. aus ganzen relativ primen Zahlen $x, y \neq 0, 0$ bestehende) Lösung der Ungleichung $|\xi\eta| < \frac{1}{2}$, und es enthält höchstens zwei verschiedene solcher Lösungen, wobei dann die Determinante aus den Koordinaten x, y dieser Lösungen stets ± 1 ist. Zwei entgegengesetzte Systeme x, y und $-x, -y$ betrachten wir hier nicht als verschieden. Umgekehrt gibt es zu jeder primitiven Lösung x, y eine solche Form jenes Tangentenparallelogramms, wobei es nur diese Lösung (und $-x, -y$) enthält, und andererseits, wofern dafür nicht $\xi = 0$ ist, auch eine solche Form, wobei es diese und außerdem noch eine zweite primitive Lösung mit kleinerem $|\xi|$ enthält. Deformieren wir nun das Tangentenparallelogramm kontinuierlich, indem wir die Tangenten an den bezüglichen Hyperbeln entlang gleiten lassen, so erhalten wir abwechselnd die rot und grün gezeichneten Formen mit einer und mit zwei primitiven Lösungen, und es tritt die Gesamtheit der primitiven Auflösungen der diophantischen Ungleichung $|\xi\eta| < \frac{1}{2}$, geordnet nach

*) Wie die Aussage hier gefaßt ist, darf a nicht gerade die Hälfte einer ungeraden ganzen Zahl sein.

**) Wir nehmen hier an, daß $\xi\eta$ nicht gerade mit der Form $\frac{1}{2}(X^2 - Y^2)$ arithmetisch äquivalent ist.

abnehmendem $|\xi|$ (und wachsendem $|\eta|$), — die zu den Formen ξ, η gehörige Diagonalkette — hervor.

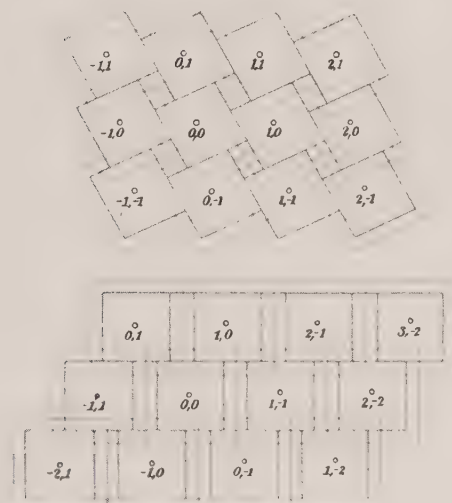
(Fig. 3). Mit den nicht homogenen diophantischen Ungleichungen hat sich wohl zuerst Tschebyscheff beschäftigt und darüber ein Theorem folgender Art entwickelt: Sind a, b zwei beliebige reelle Größen, so kann man stets ganze Zahlen x, y finden, wofür die Ungleichung (1) gilt. Statt des Faktors $\frac{1}{4}$ rechts hat Tschebyscheff hier eine etwas ungünstigere Konstante. Die am weitesten führenden Schlüsse in dieser Richtung knüpfen an die nun folgenden Zeichnungen an:

ξ, η mögen wieder zwei lineare Formen in x, y mit beliebigen reellen Koeffizienten und einer Determinante $= 1$ bedeuten. Mit Hilfe zweier aufeinanderfolgenden Glieder der vorhin besprochenen Diagonalkette zu ξ, η können wir ein System homologer Parallelogramme um die einzelnen Gitterpunkte als Mittelpunkte, die roten Parallelogramme, konstruieren, deren Diagonalen parallel den Linien $\xi = 0, \eta = 0$ sind und wobei nicht zwei ineinander eindringen, wobei aber jedes an vier (bei lückenloser Überdeckung der Ebene an sechs oder acht) andere Parallelogramme angrenzt. Dabei können sich noch zwei verschiedene Möglichkeiten darbieten, die in der oberen und der unteren Zeichnung zur Darstellung gebracht sind, indem ein Parallelogramm entweder mit jeder Seite an ein benachbartes oder nur mit zwei Seiten jedesmal an je zwei benachbarte anstößt. Die roten Parallelogramme lassen nun (im allgemeinen) noch Lücken zwischen sich. Wir dilatieren sie von ihren Mittelpunkten zu homothetischen Parallelogrammen in demjenigen bestimmten Verhältnis, wobei sie jedesmal über die Hälfte anstoßender Lücken hinauswachsen, und wir erhalten dadurch die grün gezeichneten Parallelogramme, welche nun die ganze Ebene vollständig, aber zum Teil mehrfach überdecken. Dabei zeigt es sich, daß diese neuen Bereiche keine Partie der Ebene mehr als zweifach überdecken, und ist infolgedessen der Flächeninhalt $\iint dx dy$ eines grünen Parallelogramms ≤ 2 . Diese Tatsache führt zu folgendem Satze:

Sind ξ_0, η_0 irgend zwei reelle Werte, so kann man stets solche ganze Zahlen x, y finden, daß die Ungleichung (3) gilt. Der vorhin formulierte Satz über den Ausdruck $x - ay - b$ ist nur ein Spezialfall dieser allgemeineren Aussage. —

Das arithmetische Fundamentaltheorem über konvexe Gebilde läßt sich mit Leichtigkeit auf den Raum, sowie auf Mannigfaltigkeiten beliebiger Dimension übertragen. Eine seiner wertvollsten Anwendungen findet das Theorem zu einfachen Beweisen der Dirichletschen Sätze über die Einheiten in den algebraischen Zahlkörpern.

Fig. 3. Inhomogeneous lineare Ausdrücke.

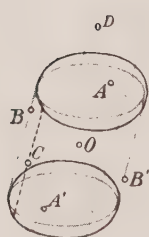


$$(1) \quad |x - ay - b| < \frac{1}{4|y|}.$$

$$(2) \quad \xi = \alpha x + \beta y, \quad \eta = \gamma x + \delta y, \\ \alpha \delta - \beta \gamma = 1;$$

$$(3) \quad |(\xi - \xi_0)(\eta - \eta_0)| < \frac{1}{4}.$$

Fig. 4. Kettenalgorithmen für drei lineare Formen.



1° Eine reelle, zwei konjugiert komplexe Formen:

$$\left\{ \begin{array}{l} \xi = \alpha x + \beta y + \gamma z, \\ \eta = (\alpha' + i\alpha'')x + (\beta' + i\beta'')y + (\gamma' + i\gamma'')z, \\ \zeta = (\alpha' - i\alpha'')x + (\beta' - i\beta'')y + (\gamma' - i\gamma'')z; \end{array} \right.$$

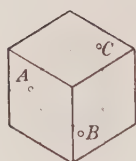
$$(1) \quad -\lambda \leq \xi \leq \lambda, \quad |\eta| \leq \mu.$$

2° Drei reelle Formen:

$$\begin{aligned} \xi &= \alpha x + \beta y + \gamma z, \\ \eta &= \alpha' x + \beta' y + \gamma' z, \quad \text{Det.} \neq 0; \\ \zeta &= \alpha'' x + \beta'' y + \gamma'' z, \end{aligned}$$

$$(2) \quad |\xi|, |\eta|, |\zeta| \leq \sqrt[3]{|\text{Det.}|}.$$

I



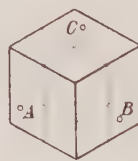
$A + B + C$
nicht im Inneren des
Parallelepipeds;

II



$A + B$ u.
 $A + B - C$
nicht im Inneren;

III



$A + B + C$
 $= 0.$

Im Raume werden wir so für das parallelepipedische Punktgitter der ganzzahligen Systeme von 3 Variablen x, y, z den Satz haben: *Ein konvexer Körper mit einem Mittelpunkt im Nullpunkt und von einem Volumen $\iiint dx dy dz \leq 8$ enthält stets wenigstens einen weiteren Gitterpunkt.*

Auf diesen Satz gründen wir wirklich brauchbare *Methoden zur Ermittlung der sämtlichen Einheiten in einem gegebenen kubischen Zahlkörper.*

(Fig. 4.) Handelt es sich zunächst um einen kubischen reellen Körper mit *negativer* Diskriminante, wobei die zwei konjugierten Körper einander konjugiert komplex sind, so existiert in dem Körper eine völlig bestimmte positive Fundamenteleinheit von möglichst kleinem Betrage > 1 . Das Verfahren zur Gewinnung dieser Einheit läßt sich an der oberen Zeichnung in Figur 4 klarmachen. Es sei ξ eine Basisform des gegebenen Körpers, η, ξ seien die konjugierten Formen in den konjugierten Körpern. Sind λ, μ positive Parameter, so geben uns die Gleichungen $\xi = \pm \lambda$ zwei Ebenen, $|\eta| = \mu$ eine elliptische Zylinderfläche. Wir können die zwei Parameter λ, μ derart bestimmen, daß der begrenzte zylindrische Raum (1) (der rote Zylinder in der Figur) sowohl auf einer Basisfläche einen Gitterpunkt A wie auf der Mantelfläche einen Gitterpunkt B (und natürlich gleichzeitig auch die in bezug auf den Nullpunkt diametral gegenüberliegenden Gitterpunkte A', B') aufweist. Nun bedürfen wir eines neuen Gitterpunktes C , um hernach aus den Koordinaten von A, B, C eine hier nützliche lineare Substitution zu entnehmen. Wir variieren den elliptischen Zylinder (1) als solchen, indem wir in den zwei Basisflächen diejenigen parallelen Durchmesser festhalten, deren Ebene durch B, B' geht, dagegen die zu ihnen konjugierten Durchmesser dilatieren und entsprechend den ganzen zylindrischen Raum ausdehnen, bis wir einen neuen Gitterpunkt C auf seiner Mantelfläche auftreten sehen. In diesem Zustande (mit den schwarz gezeichneten Rändern) bezeichnen wir den Zylinder als einen *extremen*. Wir finden, daß die Determinante aus den Koordinaten von A, B, C gleich ± 1 ist, wenn sie nicht unter besonderen Umständen Null ist. Von jedem extremen Zylinder können wir nun zu einem benachbarten schmäleren extremen Zylinder fortschreiten. Wir ziehen die ursprünglichen Basisflächen homothetisch von der Achse aus zusammen, bis ihre Ränder durch A, A' gehen, wodurch diese Punkte auf die Mantelfläche treten, und können dann, ohne daß Gitterpunkte ins Innere des Zylinders eintreten, den Zylinder ausziehen, die Basisflächen parallel mit ihrer Anfangslage vom Nullpunkte entfernen, bis sie von neuem auf Gitterpunkte D, D' stoßen, und von diesem weiteren Zustande aus können wir dann wie vorhin einen extremen Zylinder herstellen. Danach existiert eine bestimmte *Kette von extremen Zylindern*, und die

Betrachtung der auf ihren Basisflächen auftretenden Gitterpunkte führt mit Sicherheit eben zur Auffindung der Fundamenteinheit in dem gegebenen kubischen Zahlkörper. —

Das allgemeine Theorem über konvexe Körper, auf Parallelepipede angewandt, führt zur Folgerung: *Sind ξ, η, ζ irgend drei lineare Formen in x, y, z mit beliebigen reellen Koeffizienten und einer Determinante $\pm \Delta$, wobei $\Delta > 0$ ist, so kann man für die Variablen x, y, z stets solche ganze Zahlen, die nicht sämtlich Null sind, finden, daß dadurch die Beträge aller drei Formen $\leq \sqrt[3]{\Delta}$ ausfallen.*

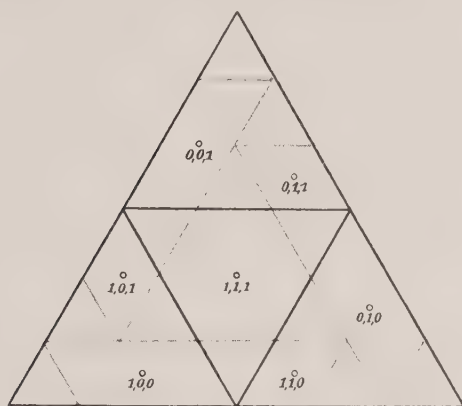
Liegt nun ein kubischer Körper mit positiver Diskriminante vor, dessen konjugierte Körper also sämtlich reell sind, und handelt es sich um die Ermittlung aller Einheiten im Körper, so seien ξ, η, ζ konjugierte Basisformen in dem Körper und seinen zwei konjugierten Körpern. Wir fassen alsdann die Gesamtheit aller solchen „extremen“ Parallelepipede mit dem Nullpunkt als Mittelpunkt und mit Seitenflächen parallel den Ebenen $\xi = 0, \eta = 0, \zeta = 0$ ins Auge, welche im Inneren vom ganzen Gitter nur den Nullpunkt enthalten, aber auf allen Seitenflächen mit besonderen Gitterpunkten versehen sind. Es existieren hier unendlich viele Parallelepipede von diesem Charakter, sie besitzen eine bestimmte Verkettung untereinander, und die Kenntnis dieser führt uns mit Sicherheit zur Aufstellung aller Einheiten im gegebenen Zahlkörper. Den Übergang von einem extremen Parallelepiped zu seinen benachbarten in der Kette vermittelt ein einfacher Algorithmus, der sich vor allem nach der Art richtet, wie die Gitterpunkte auf den Seitenflächen des Parallelepipeds in Hinsicht auf deren Mittellinien liegen. In dieser Beziehung bieten sich wesentlich drei Möglichkeiten dar, die in den Figuren (I), (II), (III) zum Ausdruck gebracht sind. In den Fällen (I) und (II) erweist sich die Determinante aus den Koordinaten von A, B, C gleich 1, im Falle (III) ist sie Null und fällt dabei der Schwerpunkt des Dreiecks ABC in den Nullpunkt.

(Fig. 5.) Ich hebe noch das folgende anziehende Problem hervor, welches auch in der Theorie von der Struktur der Kristalle eine Stelle findet:

Wir denken uns einen beliebigen Grundkörper im Raume vorgelegt. Lauter mit ihm kongruente und parallel orientierte Körper in unendlicher Anzahl seien so angeordnet, daß ihre Schwerpunkte ein parallelepipedisches Punktsystem bilden und daß nicht zwei der Körper ineinander eindringen. Unter welchen Umständen schließen sich die Körper so dicht als überhaupt möglich zusammen, sind also die zwischen ihnen vorhandenen Lücken auf ein Minimum an Volumen reduziert?

Für den Fall, daß der Grundkörper ein Oktaeder ist, gibt Fig. 5 die

Fig. 5. Dichteste Lagerung von Oktaedern.



$$\varphi = -\xi + \eta + \zeta, \quad \chi = \xi - \eta + \zeta, \quad \psi = \xi + \eta - \zeta, \\ \omega = \xi + \eta + \zeta,$$

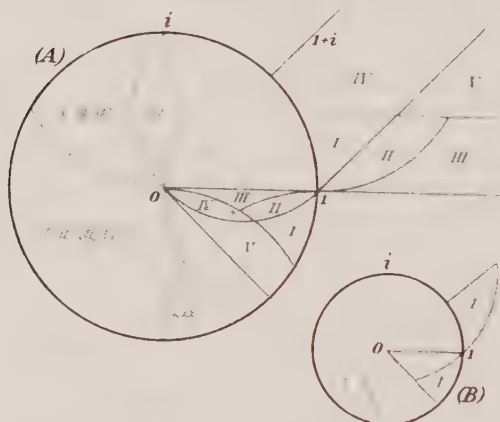
$$(1) \quad \varphi + \chi + \psi + \omega = 0, \quad \text{Det. } (\xi, \eta, \zeta) = \Delta;$$

$$(2) \quad |\varphi|, |\chi|, |\psi|, |\omega| \leq \sqrt[3]{\frac{108}{19}} \Delta.$$

$$(3) \quad \pm(x - az) \pm \frac{z}{t} = 1, \quad \pm(y - bz) \pm \frac{z}{t} = 1;$$

$$(4) \quad \left| \frac{x}{z} - a \right|, \left| \frac{y}{z} - b \right| < \sqrt[3]{\frac{8}{19} \frac{1}{z^{\frac{3}{2}}}}, \quad \left(\sqrt[3]{\frac{8}{19}} = 0,648 \dots \right)$$

Fig. 6. Lineare Formen im Körper von $i = \sqrt[4]{1}$.



$$(1) \quad \begin{aligned} \xi &= (\alpha + i\alpha')(x + ix') + (\beta + i\beta')(y + iy'), \\ \eta &= (\gamma + i\gamma')(x + ix') + (\delta + i\delta')(y + iy'), \end{aligned} \quad \text{Det.}(\xi, \eta) = \Delta;$$

$$(2) \quad \begin{vmatrix} x_1 + ix_1' & x_2 + ix_2' \\ y_1 + iy_1' & y_2 + iy_2' \end{vmatrix}, \quad (8) \quad \begin{aligned} \xi &= \lambda e^{i\varphi}(X + iX' + \varrho(Y + iY')), \\ \eta &= \mu e^{i\psi}(\sigma(X + iX') + Y + iY'), \end{aligned}$$

$$(4) \quad |\xi|, |\eta| \leq \sqrt{\frac{\sqrt{3} + 1}{\sqrt{6}}} |\Delta|$$

Lösung des Problems an. In der fraglichen dichtesten gitterförmigen Lagerung muß jedes einzelne der Oktaeder in bestimmter Weise an vierzehn benachbarte anstoßen. Hier ist, in eine Ebene umgeklappt, das halbe Netz eines der Oktaeder dargestellt und sind in den vier Seitenflächen (durch zur Hälfte rote Berandung) die 7 Partien mit Mittelpunkt angezeigt, in welchen das Oktaeder sich an benachbarte anlegt. Das Minimum des Raumes, welches die Lücken zwischen den Oktaedern noch darbieten, verhält sich zu dem von den Oktaedern eingenommenen Raume in diesem Falle der dichtesten Lagerung wie 1:19. Dieses Resultat gestattet folgende rein arithmetische Einkleidung:

Sind $\varphi, \chi, \psi, \omega$ irgend vier lineare Formen in x, y, z mit beliebigen reellen Koeffizienten, deren Summe identisch Null ist und wobei je drei eine Determinante $\pm 4\Delta$ ($\Delta > 0$) haben, so kann man stets solche ganze Zahlen x, y, z , die nicht sämtlich Null sind, finden, daß die Ungleichungen (2) gelten.

Von diesem Ergebnisse machen wir noch eine bemerkenswerte Anwendung. Es seien a, b zwei beliebige reelle Größen, t ein positiver Parameter, so bestimmen die 8 Ebenen (3) ein Oktaeder. Indem noch t beliebig groß angenommen werden kann, entspringt hieraus diese Folgerung:

Man kann zwei beliebige reelle Größen a, b stets durch Brüche x/z und y/z mit gleichem Nenner beliebig genau und zugleich derart annähern, daß die Ungleichungen (4) statthaben.

(Fig. 6.) Wir werfen nun die Frage auf: Wie übertragen sich die Sätze über die Approximation einer reellen Größe durch Zahlen des natürlichen Rationalitätsbereiches auf das Gebiet der komplexen Größen? Man wird hier zunächst auf Approximationen im Zahlkörper der dritten oder der vierten Einheitswurzel ausgehen. Im Körper der dritten Einheitswurzel liegen die Dinge wesentlich einfacher und werden die Sätze sehr ähnlich denen für reelle Größen. Ich will hier nur die verwickelteren Beziehungen im Körper der vierten Einheitswurzel berühren.

Es seien ξ, η zwei lineare Formen mit beliebigen komplexen Koeffizienten und zwei komplexen Variablen $x + ix', y + iy'$ von einer Determinante $\Delta \neq 0$, so richten wir unser Augenmerk auf diejenigen „extremen“ Zahlenpaare $x + ix', y + iy' \neq 0, 0$ im Zahlkörper von i , zu welchen nicht ein Zahlenpaar derselben Art angebbar ist, das sowohl $|\xi|$ wie $|\eta|$ kleiner werden läßt. Wir können die zwei Zahlen eines Paares mit einer und derselben Einheit $-1, \pm i$ multiplizieren, das entstehende assoziierte Paar gilt uns hier als nicht verschieden von dem ursprünglichen. Alle vorhandenen extremen Paare lassen sich nun in eine Kette nach der Größe von $|\xi|$ (und zugleich von $|\eta|$) ordnen. Zwei benachbarte Paare der Kette zusammen sind leicht a priori zu charakterisieren. Nämlich die

zugehörige Determinante (2) ist entweder (A) eine Einheit $\pm 1, \pm i$ oder (B) gleich $(1 + i)$, multipliziert in eine Einheit. Wir benutzen diese zwei Paare als Vertikalreihen der Matrix einer auf ξ, η anzuwendenden Substitution und erhalten dadurch für ξ, η die Ausdrücke (3), worin $|\varrho|, |\sigma|$ beide ≤ 1 sind. Indem wir das zweite Paar durch ein assoziiertes ersetzen und eventuell noch i in $-i$ umwandeln, können wir ϱ auf den in den Zeichnungen rot markierten Oktanten des Einheitskreises (bzw. $1/\varrho$ auf den konjugierten Oktanten außerhalb dieses Kreises) bringen. Es sind nun die Fälle (A) und (B) zu unterscheiden, auf welche sich die größere bzw. die kleinere Zeichnung bezieht. Im Falle (A) wird jener Oktant durch gewisse Kreise vom Radius 1 bzw. $\frac{1}{\sqrt{2}}$ in fünf Stücke zerlegt, die

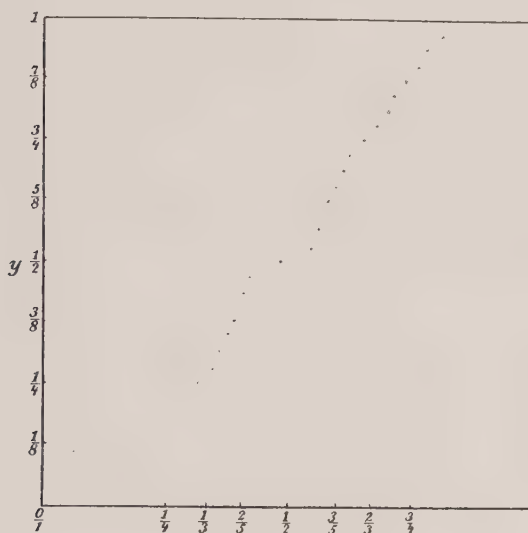
in der Figur fortlaufend mit I—V numeriert sind. Wenn ϱ in ein bestimmtes derselben fällt, kann jedesmal σ nur in diejenigen grünbegrenzten Stücke des Einheitskreises fallen, in welche die nämliche Nummer eingetragen ist. Der kleinste Wert für den Betrag der Determinante $|1 - \varrho\sigma|$ entsteht, wenn ϱ, σ den scharfen mit kleinen Kreuzen bezeichneten Ecken der Figur entsprechen. Im Falle (B) kann ϱ nur in das rote Gebiet (I) und σ dann nur in das grüne Gebiet (I) fallen. Als wichtigstes Ergebnis entnehmen wir hieraus:

Man kann in die Formen ξ, η für $x + ix', y + iy'$ stets solche ganze Zahlen des Körpers von i , die nicht beide Null sind, setzen, daß dabei die Ungleichungen (4) gelten. —

Endlich möchte ich noch einige Worte über *Kriterien für algebraische Zahlen* hinzufügen.

(Fig. 7.) Durch diese Figur suche ich dem bekannten Lagrangeschen Kriterium für eine reelle quadratische Irrationalzahl eine neue Seite abzugewinnen. In einem Quadrat von der Seitenlänge 1 sind hier auf der linken vertikalen Seite, der y -Achse, fortgesetzt Halbierungen vorgenommen, so daß sukzessive alle Punkte erhalten werden, deren Ordinate eine rein dyadische Zahl, d. h. eine rationale Zahl mit einer Potenz von 2 als Nenner ist. Jedem auf der y -Achse auftretenden Intervall oder Teilpunkt wird nun ein Intervall bzw. ein rationaler Teilpunkt auf der x -Achse, der unteren horizontalen Seite, dadurch zugeordnet, daß zunächst den Endwerten $y = 0$ und $y = 1$ die Endwerte $x = 0$ und $x = 1$ entsprechen sollen, und weiter, so oft dort ein Intervall halbiert wird, hier zwischen die Endpunkte $a/b, a'/b'$ des zugeordneten Intervalls, a und b , ferner a' und b' als relativ prim gedacht, ein neuer Teilpunkt in $x = (a + a')/(b + b')$ eingeschaltet wird. Auf der horizontalen Seite treten so als Teilpunkte sukzessive alle Punkte mit rationaler Abszisse auf, und die Zuordnung der gleichzeitig konstruierten Abszissen und Ordinaten liefert uns das Bild

Fig. 7.
Kriterium für die reellen quadratischen Irrationalzahlen.



$$x = \frac{a}{b} \dots \frac{a + a'}{b + b'} \dots \frac{a'}{b'}$$

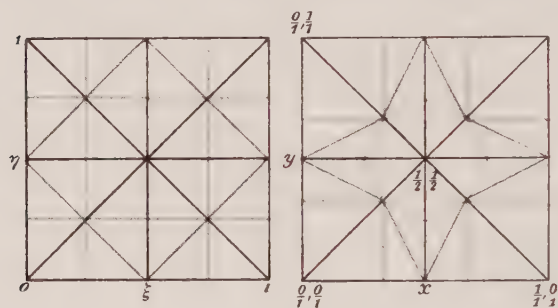
x quadratische Irrationalzahl, y rational
und nicht dyadisch;

$y = ?(x)$:

x rational,

y dyadisch.

Fig. 8.
Kriterium für die reellen kubischen Irrationalzahlen.



$$\frac{a}{c}, \frac{b}{c} \dots \frac{a+a'}{c+c'}, \frac{b+b'}{c+c'} \dots \frac{a'}{c'}, \frac{b'}{c'},$$

$$(1) \quad \xi = \varphi(x, y), \quad \eta = \psi(x, y).$$

1, x, y unabhängige Zahlen in einem kubischen Körper;
 ξ, η rational,
 von den Zahlen $\xi, \eta, \xi - \eta, \xi + \eta$ keine dyadisch.

einer beständig wachsenden Funktion $y = ?(x)$, zunächst für alle rationalen x , dann durch die Forderung der Stetigkeit erweitert auf beliebige reelle Argumente im Intervalle $0 \leq x \leq 1$, während gleichzeitig y dieses Intervall beliebig durchläuft. Wenn nun x eine quadratische Irrationalzahl ist und daher auf eine periodische Kettenbruchentwicklung führt, so entspricht dadurch dem Werte $y = ?(x)$ eine periodische Dualentwicklung und erweist sich infolgedessen y als rational. Wir erhalten dadurch die Sätze:

Ist x eine quadratische Irrationalzahl, so ist y rational, aber nicht rein dyadisch. Ist x rational, so ist y rein dyadisch. Und diese Sätze sind völlig umkehrbar.

(Fig. 8.) Die letzte Figur zeigt nun eine Verallgemeinerung dieser Sätze auf *kubische Irrationalitäten*, welche Herr Louis Kollros in seiner Dissertation (Zürich 1904) entwickelt hat:

Einerseits wird hier ein Quadrat, in welchem ξ, η in den Grenzen 0 und 1 laufen, in der Weise behandelt, daß es zuerst durch die Diagonale $\xi = \eta$ in zwei gleichschenklige rechtwinklige Dreiecke zerlegt wird und in der Folge jedes einmal entstehende gleichschenklige rechtwinklige Dreieck durch die Verbindung von der Spitze nach der Mitte der Hypotenuse in zwei gleichschenklige rechtwinklige Dreiecke aufgelöst wird. Andererseits erfährt gleichzeitig ein zweites Quadrat, in welchem x, y in den Grenzen 0 und 1 laufen, eine gewisse Schritt für Schritt zugeordnete Zerfällung in Dreiecke: zunächst werden die vier Ecken des Quadrats den vier Ecken des ersten Quadrats mit gleichwertigen Koordinaten und weiter die Linie $x = y$ der Linie $\xi = \eta$ zugeordnet, und in der Folge wird, wo dort eine Hypotenuse halbiert und hernach eine geradlinige Verbindung eingeführt wird, hier zwischen die Endpunkte der zugeordneten Strecke, wenn deren Koordinaten $a/c, b/c$ und $a'/c', b'/c'$ und a, b, c sowie a', b', c' relativ prime Zahlen sind, als ein neuer Teilpunkt der Punkt mit den Koordinaten $x = (a + a')/(c + c')$, $y = (b + b')/(c + c')$ eingeschaltet und hernach die entsprechende geradlinige Verbindung vorgenommen. Dadurch werden zwei eindeutig umkehrbare Beziehungen (1) festgesetzt, zunächst für alle rationalen x, y und dyadischen ξ, η und hernach durch die Forderung der Stetigkeit überhaupt für beliebige Argumente und Funktionswerte in den Grenzen 0 und 1. Dabei findet nun Kollros die folgenden Sätze, welche freilich in einem wesentlichen Punkte noch nicht bewiesen sind, deren Richtigkeit aber nach einer Menge von Beispielen in hohem Grade plausibel erscheint:

Sind 1, x, y drei unabhängige Zahlen in einem reellen kubischen Körper, so sind ξ, η rational und ist keine der Größen $\xi, \eta, \xi + \eta, \xi - \eta$ rein dyadisch. Gehören x, y einem quadratischen Körper an, ohne beide rational

zu sein, so sind ξ, η rational und ist eine der Größen $\xi, \eta, \xi + \eta, \xi - \eta$ rein dyadisch. Sind x, y beide rational, so sind ξ, η beide rein dyadisch. Diese Sätze sind vollkommen umkehrbar.

Nimmt man $y = x^2$, so erlangt man hierdurch ein vollständiges Kriterium dafür, daß x eine reelle kubische Irrationalzahl ist. Besonders zu betonen ist, daß diese Sätze für alle kubischen Körper und nicht etwa bloß für solche mit negativer Diskriminante, in welchen nur eine Fundamenteinheit vorhanden ist, zuzutreffen scheinen.

Diese Aufzählung von speziellen Ergebnissen aus der Geometrie der Zahlen ließe sich noch weiterführen. Aber ich habe vielleicht mein Ziel bereits erreicht und Sie mögen den Eindruck gewonnen haben, daß es sich hier um Fragen handelt, welche die Fundamente der Größenlehre berühren, welche der Auffassung leicht zugänglich sind und welche uns die Disziplinen der Algebra, Arithmetik, Geometrie in harmonischer Wechselwirkung zeigen.

XXI.

Diskontinuitätsbereich für arithmetische Äquivalenz.

(Crelles Journal für die reine und angewandte Mathematik, Band 129, S. 220—274.)

Problemstellung.

Die vorliegende Arbeit benutzt mehrfach Methoden, welche von Dirichlet ausgebildet worden sind.

Ich stellte mir folgendes Problem:

Ein System von n linearen Formen $\xi_1, \xi_2, \dots, \xi_n$ mit n Variablen $x_1, x_2, \dots, x_n$, mit beliebigen reellen Koeffizienten

$$\frac{\partial \xi_h}{\partial x_k} = \alpha_{hk}$$

und von Null verschiedener Determinante heiße einem zweiten solchen Systeme $\eta_1, \eta_2, \dots, \eta_n$ *arithmetisch äquivalent*, wenn jedes der zwei Systeme sich in das andere durch eine homogene lineare *Substitution mit lauter ganzzahligen Koeffizienten* transformieren läßt. Es soll nun *in der Mannigfaltigkeit A der n^2 reellen Parameter α_{hk} ein Bereich B* konstruiert werden, in welchem *jede vollständige Vereinigung (Klasse) untereinander äquivalenter Systeme durch einen Punkt*, und wenn der Punkt ins Innere des Bereichs zu liegen kommt, auch nur durch einen einzigen Punkt repräsentiert wird.

Dieses Problem läßt sich sofort auf ein entsprechendes Problem über positive quadratische Formen zurückführen. Wir bilden aus $\xi_1, \xi_2, \dots, \xi_n$ die Quadratsumme

$$\xi_1^2 + \xi_2^2 + \dots + \xi_n^2 = f.$$

Sie ist eine *positive quadratische Form der n Variablen $x_1, x_2, \dots, x_n$ mit den Koeffizienten*

$$\frac{1}{2} \frac{\partial^2 f}{\partial x_h \partial x_k} = \alpha_{hk} = \alpha_{1h} \alpha_{1k} + \alpha_{2h} \alpha_{2k} + \dots + \alpha_{nh} \alpha_{nk}.$$

Wir betrachten die $\frac{n(n+1)}{2}$ -dimensionale Mannigfaltigkeit A der $\frac{n(n+1)}{2}$ beliebig veränderlichen reellen Parameter α_{hk} . Jedem Punkte $f = (\alpha_{hk})$ dieser Mannigfaltigkeit A , für den die Form f sich als wesentlich positiv

erweist, entspricht auf Grund der eben hingeschriebenen Gleichungen noch ein $\frac{n(n-1)}{2}$ -dimensionales Gebiet $A(f)$ von Punkten (α_{hk}) in der Mannigfaltigkeit A .

Wir übertragen den Begriff der arithmetischen Äquivalenz und Klasse sinngemäß auf die positiven quadratischen Formen, und *wir suchen in der Mannigfaltigkeit A einen Bereich B , in dem jede Klasse positiver quadratischer Formen durch einen Punkt, und wenn der Punkt in das Innere von B fällt, auch nur durch einen einzigen Punkt repräsentiert wird.*

Alsdann bilden die Gebiete $A(f)$ zu allen in B gelegenen Punkten f den verlangten Diskontinuitätsbereich B in der Mannigfaltigkeit A .

Ich werde nachweisen, daß *der Diskontinuitätsbereich B für die arithmetische Äquivalenz der positiven quadratischen Formen* derart konstruiert werden kann, daß *er in der Mannigfaltigkeit A einen von einer endlichen Anzahl von Ebenen begrenzten konvexen Kegel mit dem durch $f=0$ repräsentierten Nullpunkte als Spitze vorstellt.*

Durch diesen Satz wird für die positiven quadratischen Formen mit n Variablen die Theorie der arithmetischen Äquivalenz auf dieselbe Höhe gebracht, welche sie für die ternären Formen durch die Abhandlung von Dirichlet: *Über die Reduction der positiven quadratischen Formen mit drei unbestimmten ganzen Zahlen* erreichte.

Derjenige Teil des Diskontinuitätsbereichs B , welcher den Formen f mit einer Determinante ≤ 1 entspricht, besitzt ein endliches Volumen. Für $n=2$ kommt dieses Volumen wesentlich auf den nichteuklidischen Flächeninhalt des Fundamentalbereichs der elliptischen Modulfunktion $J(\omega)$ hinaus. Die allgemeine Ermittlung jenes Volumens gelingt hier mit Hilfe derjenigen Prinzipien, auf welche Dirichlet die Berechnung der Klassenanzahlen der ganzzahligen binären quadratischen Formen gegründet hat, und zwar kommt das analytische Element in diesen Prinzipien gerade bei der hier zu machenden Anwendung am reinsten zum Vorschein.

Mit jenem Volumen hängen gewisse asymptotische Gesetze betreffs der Formen mit ganzzahligen Koeffizienten zusammen. Andererseits ermöglicht der Wert des Volumens einen Schluß auf die dichteste Ausfüllung des n -dimensionalen Raumes durch lauter kongruente Kugeln.

Die einschlägige Literatur zu den hier behandelten Gegenständen ist in meiner Arbeit im 107. Bande von Crelles Journal (Diese Ges. Abhandlungen, Bd. I, S. 243—260) ausführlich genannt und verweise ich dieserhalb auf die dortigen Angaben.

§ 1. Charakter der positiven quadratischen Formen.

Es sei

$$f(x_1, x_2, \dots, x_n) = \sum a_{hk} x_h x_k \quad (h, k = 1, 2, \dots, n)$$

eine quadratische Form der n Variablen $x_1, x_2, \dots, x_n$ mit reellen Koeffizienten a_{hk} . Wir setzen stets $a_{kh} = a_{hk}$ für $h < k$ voraus. Die aus der $h_1, h_2, \dots, h_\nu$ -ten Horizontal- und der $k_1, k_2, \dots, k_\nu$ -ten Vertikalreihe des quadratischen Schemas der a_{hk} hergestellte ν -reihige Determinante bezeichnen wir mit

$$D \begin{pmatrix} h_1, h_2, \dots, h_\nu \\ k_1, k_2, \dots, k_\nu \end{pmatrix}.$$

1. Dafür, daß f eine *wesentlich positive* Form sei, ist bekanntlich notwendig und hinreichend, daß die n Determinanten

$$D \begin{pmatrix} 1, 2, \dots, h \\ 1, 2, \dots, h \end{pmatrix} = D_h \quad \text{für } h = 1, 2, \dots, n$$

sämtlich positiv sind. Die letzte dieser Größen, D_n , ist die Determinante der Form f , deren Wert wir auch mit $D(f)$ bezeichnen. Wir setzen noch $D_0 = 1$,

$$\frac{D_h}{D_{h-1}} = q_h \quad (h = 1, 2, \dots, n); \quad \frac{D \begin{pmatrix} 1, \dots, h-1, h \\ 1, \dots, h-1, k \end{pmatrix}}{D \begin{pmatrix} 1, \dots, h-1, h \\ 1, \dots, h-1, h \end{pmatrix}} = \gamma_{hk} \quad \begin{pmatrix} h = 1, \dots, n-1 \\ k = h+1, \dots, n \end{pmatrix}.$$

Dann sind auch alle Größen $q_h > 0$ und besteht die identische Darstellung

$$(1) \quad f = q_1 \xi_1^2 + q_2 \xi_2^2 + \dots + q_n \xi_n^2$$

mit

$$(2) \quad \xi_1 = x_1 + \gamma_{12} x_2 + \dots + \gamma_{1n} x_n, \quad \xi_2 = x_2 + \dots + \gamma_{2n} x_n, \dots, \xi_n = x_n,$$

welche den *Charakter von f als positive Form in Evidenz* setzt.

2. Die Vergleichung der Koeffizienten von $x_1^2, x_2^2, \dots, x_n^2$ auf beiden Seiten von (1) ergibt

$$(3) \quad a_{11} = q_1, \quad a_{22} \geq q_2, \dots, a_{nn} \geq q_n$$

und durch Multiplikation dieser Ungleichungen folgt

$$(4) \quad a_{11} a_{22} \dots a_{nn} \geq D(f).$$

3. Ist L ein gegebener positiver Wert und soll $f \leq L$ ausfallen, so folgt aus (1):

$$|\xi_n| \leq \sqrt{\frac{L}{q_n}}, \quad |\xi_{n-1}| \leq \sqrt{\frac{L}{q_{n-1}}}, \dots, |\xi_1| \leq \sqrt{\frac{L}{q_1}},$$

und diesen Ungleichungen können nach den Ausdrücken (2) nur eine

endliche Anzahl verschiedener ganzzahliger Systeme $x_n, x_{n-1}, \dots, x_1$ genügen. Wir haben daher den Satz:

Eine positive quadratische Form f kann nur für eine endliche Anzahl von ganzzahligen Systemen der Variablen Werte annehmen, die eine gegebene Schranke L nicht überschreiten.

§ 2. Anordnung in einer n -dimensionalen Mannigfaltigkeit.

Sind $a_1, a_2, \dots, a_n$ und $b_1, b_2, \dots, b_n$ zwei verschiedene Systeme von je n Größen und hat man

$$a_1 = b_1, \dots, a_{l-1} = b_{l-1}, \quad a_l > b_l,$$

wo l einen der Werte $1, 2, \dots, n$ bedeuten kann, so heiße das erste System *höher*, das zweite *niedriger* als das andere, und zwar an l^{ter} Stelle höher bzw. niedriger.

Es sei eine unendliche Menge von Systemen $S(a_1, a_2, \dots, a_n)$ vorgelegt mit folgender Eigenschaft: Ist $a_1, a_2, \dots, a_n$ ein beliebiges System daraus und l eine beliebige der Zahlen $1, 2, \dots, n$, so soll bei allen vorhandenen Systemen $b_1, b_2, \dots, b_n$ der Menge, welche genau an l^{ter} Stelle niedriger als das erste System sind, für die Größe b_l jedesmal nur eine *endliche* Anzahl verschiedener Werte in Betracht kommen.

Bildet man alsdann, von einem beliebigen Systeme S der Menge ausgehend, soweit als zugänglich, eine Reihe von Systemen $S, S^{(1)}, S^{(2)}, \dots$, so daß jedes folgende niedriger als das vorhergehende ist, so muß eine solche Reihe stets nach einer endlichen Anzahl von Schritten abbrechen. Es muß sich also nach einer endlichen Anzahl von Schritten schließlich ein System einstellen, zu welchem es kein niedrigeres in der Menge gibt, und welches demnach das *niedrigste* System in der Menge vorstellt.

Für $n = 1$ ist die Behauptung evident. Nun sei $n > 1$. Betrachten wir zum Ausgangssysteme $S(a_1, a_2, \dots, a_n)$ alle Systeme $a_1, b_2, \dots, b_n$ der Menge, welche mit ihm in dem ersten Elemente a_1 übereinstimmen, so bilden die darin auftretenden Systeme $b_2, \dots, b_n$ eine Menge von Systemen aus $n - 1$ Elementen mit ganz entsprechender Eigenschaft, wie wir sie für die gegebene Menge n -gliedriger Systeme voraussetzen. Nehmen wir nun den zu beweisenden Satz bereits als erwiesen an, wenn die Zahl $n - 1$ anstatt n steht, so kann eine Reihe $S, S^{(1)}, S^{(2)}, \dots$ nur mit einer *endlichen* Anzahl von Termen derart fortschreiten, daß das erste Element ungeändert bleibt und jedesmal eine *Erniedrigung* an der zweiten bis n^{ten} Stelle eintritt. Da außerdem eine *Erniedrigung* genau an erster Stelle nach Voraussetzung ebenfalls nur eine endliche Anzahl von Malen vorkommen kann, so ist der zu beweisende Satz sofort für Systeme aus n Gliedern klar.

§ 3. Niedrigste Formen in einer Klasse.

Der Begriff der arithmetischen Äquivalenz von positiven quadratischen Formen ist schon in der Einleitung erwähnt worden. Zwei positive Formen

$$f = \sum a_{hk} x_h x_k, \quad g = \sum b_{hk} y_h y_k \quad (h, k = 1, \dots, n)$$

sind hierbei dann und nur dann *äquivalent*, wenn f in g durch eine *ganzzahlige Transformation mit einer Determinante ± 1* überzuführen ist.

1. Ist dabei

$$(5) \quad a_{11} = b_{11}, \quad a_{22} = b_{22}, \dots, a_{nn} = b_{nn},$$

so mögen f und g *gleichgestellt* heißen. Ist dagegen

$$(6) \quad a_{11} = b_{11}, \dots, a_{l-1, l-1} = b_{l-1, l-1}, \quad a_{ll} > b_{ll},$$

wobei l einen der Werte $1, 2, \dots, n$ haben kann, so soll f *höher* als g und g *niedriger* als f und zwar an l^{ter} Stelle höher bzw. niedriger heißen.

Es sei

$$(7) \quad x_h = s_{h1}y_1 + s_{h2}y_2 + \dots + s_{hn}y_n \quad (h = 1, 2, \dots, n)$$

eine ganzzahlige Substitution mit der Determinante ± 1 , welche f in g überführt, so hat man

$$b_{kk} = f(s_{1k}, \dots, s_{nk}) \quad (k = 1, 2, \dots, n).$$

Ist g an l^{ter} Stelle *niedriger* als f , so folgt daher

$$(8) \quad f(s_{11}, \dots, s_{n1}) = a_{11}, \dots, f(s_{1, l-1}, \dots, s_{n, l-1}) = a_{l-1, l-1}, \quad f(s_{1l}, \dots, s_{nl}) < a_{ll}.$$

Ferner ist die aus den l ersten Vertikalreihen der Substitution (7) gebildete *Matrix*

$$(9) \quad \|s_{hk}\| \quad (h = 1, 2, \dots, n; k = 1, 2, \dots, l)$$

unimodular, d. h. der größte gemeinsame Teiler aller aus ihr zu bildenden l -reihigen Determinanten ist gleich 1.

Hiernach ist leicht zu entscheiden, ob es in der Klasse von f eine Form g gibt, welche niedriger als f ist. Wir nehmen für l nacheinander jeden der Werte $1, 2, \dots, n$ und fassen jedesmal das System der Bedingungen (8) ins Auge. Nach § 1 kann es jedesmal gewiß nur eine *endliche* Anzahl von ganzen Zahlen s_{hk} ($k \leq l$) geben, welche diesen Bedingungen genügen. Sowie für eines der betreffenden Lösungssysteme die *Matrix* (9) *unimodular* ist, können wir bekanntlich zu dieser Matrix $n - l$ weitere Vertikalreihen ganzzahliger Koeffizienten s_{hk} ($k = l + 1, \dots, n$) hinzufügen, so daß die Determinante des entstehenden quadratischen Schemas ± 1 wird, und es führt dann die zugehörige Substitution (7) in der Tat f in eine an l^{ter} Stelle niedrigere Form g über. Dabei kommen jedesmal für b_{ll} nur eine endliche Anzahl verschiedener Werte in Frage.

2. Auf Grund des Hilfssatzes in § 2 kann man nunmehr in jeder Klasse f eine solche Form g ermitteln, zu welcher es keine niedrigere Form

in der Klasse gibt und welche wir daher eine *niedrigste* Form der Klasse nennen. Alle äquivalenten mit ihr gleichgestellten Formen sind gleichzeitig niedrigste Formen der Klasse.

3. Soll die Form f durch die Substitution (7) in eine *gleichgestellte* Form übergehen, so müssen die n Gleichungen

$$f(s_{11}, \dots, s_{n1}) = a_{11}, \dots, f(s_{1n}, \dots, s_{nn}) = a_{nn}$$

statthaben; es genügen ihnen nur eine endliche Zahl ganzzahliger Systeme s_{hk} und gibt es daher gewiß auch nur eine endliche Anzahl ganzzahliger Substitutionen mit einer Determinante ± 1 , durch welche eine Form in gleichgestellte übergeht. Insbesondere zeigt sich, daß jede positive Form nur eine endliche Anzahl ganzzahliger Transformationen in sich besitzt.

4. Jede Form f geht durch die 2^n Substitutionen

$$(10) \quad x_1 = \pm y_1, \quad x_2 = \pm y_2, \dots, x_n = \pm y_n,$$

wobei ein jedes der n Vorzeichen unabhängig von den anderen als $+$ oder $-$ angenommen werden kann, in gleichgestellte Formen über.

5. Es möge für einen Moment eine Form f und die ganze zugehörige Klasse *allgemein* heißen, wenn eine Gleichung

$$f(x_1, x_2, \dots, x_n) = f(y_1, y_2, \dots, y_n)$$

mit ganzzahligen Werten $x_1, x_2, \dots, x_n; y_1, y_2, \dots, y_n$ niemals anders bestehen kann, als daß das System $y_1, y_2, \dots, y_n$ mit $x_1, x_2, \dots, x_n$ oder mit $-x_1, -x_2, \dots, -x_n$ übereinstimmt. Insbesondere wird dieser Charakter für f zutreffen, wenn zwischen den $\frac{n(n+1)}{2}$ Koeffizienten a_{hk} von f keine homogene lineare Relation mit ganzzahligen Koeffizienten besteht. Vergleichen wir insbesondere die zwei Systeme

$$x_h = 1, x_{h+1} = 1, x_k = 0 \quad \text{und} \quad y_h = 1, y_{h+1} = -1, y_k = 0 \quad (k \neq h, h+1)$$

miteinander, so zeigt sich, daß in einer allgemeinen Form f gewiß jeder Koeffizient $a_{h,h+1}$ von Null verschieden ausfällt.

Ist die Klasse f eine allgemeine, so geht offenbar jede Form daraus einzig durch die 2^n Substitutionen (10) in äquivalente gleichgestellte Formen über, und gibt es in der Klasse stets eine *einzige niedrigste Form* g , welche noch die Bedingungen

$$b_{12} > 0, \quad b_{23} > 0, \dots, b_{n-1,n} > 0$$

erfüllt.

6. Wir bezeichnen mit E die *identische Substitution*, ferner zu jeder Substitution S als *entgegengesetzte* und mit $-S$ diejenige Substitution, welche durch Änderung aller Koeffizienten s_{hk} in die entgegengesetzten Werte $-s_{hk}$ entsteht.

§ 4. Reduzierte Formen.

Es sind wesentlich die niedrigsten Formen einer Klasse, welche Hermite (Crelles Journal, Bd. 40; Oeuvres, T. I, p. 100) als reduzierte Formen eingeführt hat mit der Absicht, in der Menge dieser Formen einen Diskontinuitätsbereich B von der in der Einleitung dargelegten Natur zu erhalten. Hier bringen wir gegenüber der Hermiteschen Definition eine Vereinfachung dadurch an, daß wir gewisse kompliziertere Charaktere, welche für niedrigste Formen zu fordern wären, die ihren Einfluß aber nur auf Teile der Begrenzung des Diskontinuitätsbereiches äußern würden, beiseite lassen.

1. Es sei l eine der Zahlen $1, 2, \dots, n$, und wir wollen unter $s_1^{(l)}, s_2^{(l)}, \dots, s_n^{(l)}$ hier *allgemein ein solches System von n ganzen Zahlen* verstehen, *wobei der größte gemeinsame Teiler der letzten $n - (l - 1)$ Zahlen darunter, also von $s_l^{(l)}, s_{l+1}^{(l)}, \dots, s_n^{(l)}$, gleich 1 ist.* Ferner bedeute $e_1^{(l)}, e_2^{(l)}, \dots, e_n^{(l)}$ speziell die l^{te} Vertikalreihe der identischen Substitution E .

Zu jedem Systeme $s_1^{(l)}, s_2^{(l)}, \dots, s_n^{(l)}$ kann man offenbar stets eine ganzzahlige Substitution $S^{(l)}$ von einer Determinante ± 1 herstellen, in deren quadratischem Koeffizientenschema die ersten $l - 1$ Vertikalreihen wie bei der identischen Substitution E aussehen und die l^{te} Reihe von eben jenem Systeme gebildet wird, also eine Substitution

$$x_h = y_h + s_h^{(l)} y_l + s_{h,l+1} y_{l+1} + \dots + s_{h,n} y_n \quad (h < l),$$

$$x_h = s_h^{(l)} y_l + s_{h,l+1} y_{l+1} + \dots + s_{h,n} y_n \quad (h \geq l).$$

Durch eine solche Substitution $S^{(l)}$ geht eine Form

$$f = \sum a_{hk} x_h x_k \quad \text{in} \quad g = \sum b_{hk} y_h y_k$$

über, so daß

$$b_{11} = a_{11}, \dots, b_{l-1,l-1} = a_{l-1,l-1}, \quad b_{ii} = f(s_1^{(l)}, s_2^{(l)}, \dots, s_n^{(l)})$$

ist. Wenn nun f speziell eine niedrigste Form in ihrer Klasse ist, wird hierbei stets $b_{ii} \geq a_{ii}$ sein müssen.

2. Wir stellen nunmehr folgende Definition auf:

Eine quadratische Form

$$f(x_1, x_2, \dots, x_n) = \sum a_{hk} x_h x_k$$

soll eine reduzierte Form heißen, wenn sie allen möglichen Ungleichungen

$$(I) \quad f(s_1^{(l)}, s_2^{(l)}, \dots, s_n^{(l)}) \geq a_{ll}$$

genügt für jedes $l = 1, 2, \dots, n$ und alle ganzzahligen Systeme $s_1^{(l)}, s_2^{(l)}, \dots, s_n^{(l)}$, wobei der größte gemeinsame Teiler der Zahlen $s_l^{(l)}, s_{l+1}^{(l)}, \dots, s_n^{(l)}$ gleich 1 ist, und ferner noch den Bedingungen:

$$(II) \quad a_{12} \geq 0, \quad a_{23} \geq 0, \dots, a_{n-1,n} \geq 0.$$

Wir schließen bei den Ungleichungen (I) jedesmal ausdrücklich die zwei Systeme $s_h^{(l)} = e_h^{(l)}$ ($h = 1, \dots, n$) und $s_h^{(l)} = -e_h^{(l)}$ ($h = 1, \dots, n$) aus, wofür (I) keine wirkliche Bedingung in den a_{hk} vorstellen würde.

Bei dieser Definition einer reduzierten Form setzen wir *nicht bereits* f als eine *positive Form* voraus.

3. In einer Klasse positiver Formen ist eine niedrigste Form, welche noch die Zusatzbedingungen (II) erfüllt, stets eine reduzierte Form. Nach den Ergebnissen des § 3 *gibt es daher in jeder Klasse positiver Formen stets wenigstens eine reduzierte Form.*

4. Die zum Index l gehörigen der Ungleichungen (I) besagen, daß a_{li} das Minimum unter allen Werten $f(x_1, x_2, \dots, x_n)$ für solche ganze Zahlen $x_1, x_2, \dots, x_n$ ist, wobei $x_i, x_{i+1}, \dots, x_n$ keinen gemeinsamen Teiler > 1 haben.

Durch eine Substitution $S^{(l)}$ gehen die sämtlichen ganzzahligen Systeme $x_1, x_2, \dots, x_n$, wobei $x_i, x_{i+1}, \dots, x_n$ keinen gemeinsamen Teiler > 1 haben, in die sämtlichen ganzzahligen Systeme $y_1, y_2, \dots, y_n$ über, wobei $y_i, y_{i+1}, \dots, y_n$ keinen gemeinsamen Teiler > 1 haben.

5. Genügt eine positive Form f den Bedingungen (I) nur für $l = 1, 2, \dots, m-1$, aber nicht für $l = m$, so können wir aus ihr durch eine Substitution $S^{(m)}$ eine äquivalente Form herleiten, welche den entsprechenden Bedingungen für $l = 1, 2, \dots, m-1$ und auch noch für $l = m$ genügt.

Man hat einfach alle ganzzahligen Systeme $s_1^{(m)}, s_2^{(m)}, \dots, s_n^{(m)}$ zu bestimmen, für welche $f(s_1^{(m)}, s_2^{(m)}, \dots, s_n^{(m)}) < a_{mm}$ ist und zugleich $s_m^{(m)}, s_{m+1}^{(m)}, \dots, s_n^{(m)}$ keinen Teiler > 1 haben. Nach Voraussetzung gibt es solche Systeme. Sie sind nur in endlicher Anzahl vorhanden. Man suche unter ihnen diejenigen heraus, welche der Form f den kleinsten Wert erteilen, und setze dann in $S^{(m)}$ als m^{te} Vertikalreihe ein beliebiges dieser Systeme an, so wird das gewünschte Ziel erreicht sein.

Beachtet man, daß der Typus $S^{(l)}$ eine beliebige ganzzahlige Substitution mit der Determinante ± 1 vorstellt, ferner, daß bei zwei verschiedenen Substitutionen $S^{(l)}$ ($l < n$) mit genau gleicher l^{ter} Vertikalreihe die zweite gleich dem Produkt der ersten in eine Substitution $S^{(l+1)}$ ist, so enthalten diese Bemerkungen eine Methode, um zu einer gegebenen positiven Form f alle ganzzahligen Substitutionen zu finden, durch welche sie in äquivalente reduzierte Formen übergeht.

Zugleich zeigt sich, daß die Anzahl dieser reduzierenden Substitutionen für jede gegebene positive Form f stets eine endliche ist.

6. Aus diesen Erwägungen leuchtet ferner ein: Eine solche reduzierte Form, für welche in keiner einzigen von den Ungleichungen (I) und (II) das Gleichheitszeichen eintritt, kann nur bei Anwendung der identischen Substitution E oder der dazu entgegengesetzten $-E$ reduziert bleiben.

7. Wir fassen nun die Koeffizienten a_{hk} einer quadratischen Form als *Koordinaten eines Punktes f in einer $\frac{n(n+1)}{2}$ -fachen Mannigfaltigkeit A* auf. In dieser Mannigfaltigkeit bezeichnen wir mit B das Gebiet derjenigen Punkte f , welche allen Ungleichungen (I) und (II) genügen. Wir nennen B den *reduzierten Raum*. Das letzte Ergebnis besagt dann:

Eine Form f im Inneren des reduzierten Raumes B , d. h. für welche in keiner der Ungleichungen (I) und (II) das Gleichheitszeichen statthat, geht durch jede von E und $-E$ verschiedene unimodulare ganzzahlige Substitution in eine Form außerhalb B über.

Insbesondere wird eine *allgemeine Klasse* immer durch einen *inneren Punkt* von B repräsentiert.

Der Bereich B geht durch jedes Paar entgegengesetzter unimodularer ganzzahliger Substitutionen $S, -S$ in eine bestimmte *äquivalente Kammer* $B_S = B_{-S}$ über, und die *Kammern*, die *verschiedenen* solchen Paaren $S, -S$ entsprechen, stoßen *untereinander höchstens in Punkten* der *Grenzungen* zusammen. Die sämtlichen Kammern B_S überdecken den ganzen Raum der positiven Formen.

§ 5. Die Wände des reduzierten Raumes.

Die Ungleichungen (I) zur Definition einer reduzierten Form f sind in *unendlicher Anzahl* vorhanden. Wir werden nunmehr den Satz beweisen:

Es gibt unter den Ungleichungen (I) eine endliche Anzahl, welche alle übrigen dieser Ungleichungen zur Folge haben.

Beweis: 1. Es sei zunächst $n = 2$, also

$$f = a_{11}x_1^2 + 2a_{12}x_1x_2 + a_{22}x_2^2.$$

Nehmen wir in (I): $s_1^{(2)} = \pm 1$, $s_2^{(2)} = 1$, so folgt

$$(11) \quad a_{11} \pm 2a_{12} + a_{22} \geq a_{22}, \quad \mp a_{12} \leq \frac{1}{2} a_{11}.$$

Aus diesen zwei Bedingungen für die zwei Vorzeichen $\pm$ geht noch $a_{11} \geq 0$ hervor. Nehmen wir $s_1^{(1)} = 0$, $s_2^{(1)} = 1$, so entsteht

$$(12) \quad a_{22} \geq a_{11}.$$

Ist $a_{11} = 0$, so folgt aus (11) auch $a_{12} = 0$ und gelten wegen $a_{22} \geq 0$ bereits alle Ungleichungen (I).

Ist $a_{11} > 0$, so folgt aus (11) und (12):

$$a_{12}^2 \leq \frac{1}{4} a_{11}^2 \leq \frac{1}{4} a_{11} a_{22}, \quad D_2 = a_{11} a_{22} - a_{12}^2 \geq \frac{3}{4} a_{11} a_{22}.$$

Schreiben wir

$$f = q_1(x_1 + \gamma_{12}x_2)^2 + q_2x_2^2,$$

so ist

$$q_1 = a_{11}, \gamma_{12} = \frac{a_{12}}{a_{11}}, \quad \pm \gamma_{12} \leq \frac{1}{2}, \quad q_2 = \frac{D_2}{q_1} \geq \frac{3}{4} a_{22}.$$

Für ganze Zahlen s_1, s_2 wird nun, wenn $|s_1| > 1, s_2 = 1$ ist:

$$f(\pm |s_1|, 1) = a_{11}(s_1^2 - |s_1|) + (a_{11} \pm 2a_{12})|s_1| + a_{22} > a_{22},$$

andererseits, wenn $|s_2| > 1$ ist:

$$f(s_1, s_2) \geq q_2 s_2^2 \geq 3a_{22} > a_{22}.$$

So leuchtet ein, daß aus (11) und (12) bereits alle Ungleichungen (I) folgen.

2. Es sei jetzt $n > 2$. Setzen wir einen Teil der Variablen $x_1, \dots, x_n$ Null und sind die noch übrigen dieser Variablen $x_{h_1}, x_{h_2}, \dots, x_{h_m} (h_1 < h_2 < \dots < h_m)$, so wird aus f eine Form $f\{x_{h_1}, x_{h_2}, \dots, x_{h_m}\}$ dieser m Variablen, und die zu (I) entsprechenden Bedingungen für diese Form sind, wie man leicht erkennt, spezielle von den Bedingungen (I) für die Form f selbst. Jede aus einer reduzierten Form f in solcher Weise abzuleitende Form von weniger als n Variablen trägt also, von den Zusatzforderungen (II) abgesehen, ebenfalls den Charakter einer reduzierten Form bei ihrer Variablenzahl.

Greifen wir nun von den Ungleichungen (I) für f erstens diejenigen heraus, welche gerade nach sich ziehen, daß die sämtlichen binären Formen

$$a_{hh}x_h^2 + 2a_{hk}x_hx_k + a_{kk}x_k^2 \quad (h < k)$$

die Bedingungen (I) für reduzierte Formen erfüllen, so handelt es sich um gewisse Ungleichungen in endlicher Anzahl und kommen diese Ungleichungen auf die folgenden

$$(13) \quad \pm 2a_{hk} \leq a_{hh} \quad (h < k)$$

und

$$(14) \quad 0 \leq a_{11} \leq a_{22} \leq \dots \leq a_{nn}$$

hinaus.

3. Wir nehmen nun an, der zu beweisende Satz sei bereits für Formen von weniger als n Variablen sichergestellt, und wir können unter den Ungleichungen (I) zweitens eine endliche Anzahl auswählen, welche zur Folge haben, daß die sämtlichen aus f abgeleiteten Formen $f\{x_{m+1}, x_{m+2}, \dots, x_n\}$ für $m = 1, 2, \dots, n-2$, reduziert sind.

Ist nun etwa

$$0 = a_{11} = a_{22} = \dots = a_{mm} \quad (1 \leq m < n), \quad 0 < a_{m+1, m+1},$$

so sind infolge von (13) überhaupt alle Koeffizienten $a_{hk} (h \leq k)$, bei welchen der Index h einen der Werte $1, 2, \dots, m$ hat, Null und wird aus f einfach $f\{x_{m+1}, \dots, x_n\}$. Die Ungleichungen (I) in bezug auf letztere Form haben dann bereits sämtliche Ungleichungen (I) für f zur Folge.

4. Nunmehr setzen wir weiterhin $a_{11} > 0$ voraus. Wir greifen drittens aus den Ungleichungen (I) diejenigen in endlicher Anzahl vorhandenen

heraus, welche zur Folge haben, daß auch *die sämtlichen aus f abzuleitenden Formen $f\{x_1, x_2, \dots, x_m\}$ für $m = 2, 3, \dots, n-1$ reduziert sind.*

5. Jetzt werden wir (in 5.—8.) zeigen, daß wir zu den bereits gewählten der Ungleichungen (I) *viertens* eine weitere *endliche* Anzahl jener Ungleichungen derart hinzufügen können, daß *aus diesen für die Determinante D_n der Form f eine Ungleichung*

$$(15) \quad D_n \geq \lambda_n a_{11} a_{22} \dots a_{nn}$$

folgt, wo λ_n eine gewisse positive nur von n und nicht von den Koeffizienten der speziellen Form f abhängende Konstante bedeutet.

Nach 1. ist diese Tatsache bereits für $n=2$ mit dem Werte $\lambda_2 = \frac{3}{4}$ zutreffend, und wir nehmen sie als *bereits erwiesen für Formen mit weniger als n Variablen* an.

Es sei jetzt m einer der Werte $1, 2, \dots, n-1$, so haben wir unter Verwendung der in § 1 eingeführten Bezeichnungen für die Determinante der Form $f\{x_1, x_2, \dots, x_m\}$ die Ungleichung:

$$(16) \quad D \begin{pmatrix} 1, 2, \dots, m \\ 1, 2, \dots, m \end{pmatrix} = D_m \geq \lambda_m a_{11} a_{22} \dots a_{mm} \quad (m \leq n-1)$$

und für die Determinante der Form $f\{x_{m+1}, x_{m+2}, \dots, x_n\}$:

$$(17) \quad D \begin{pmatrix} m+1, m+2, \dots, n \\ m+1, m+2, \dots, n \end{pmatrix} = \overline{D}_{n-m} \geq \lambda_{n-m} a_{m+1, m+1} a_{m+2, m+2} \dots a_{nn}.$$

In den Fällen $m=1$ und $m=n-1$ setzen wir $\lambda_1 = 1$.

Entwickeln wir nun den Ausdruck der Determinante D_n von f , so kommen darin $m!(n-m)!$ Terme vor, die sich zu dem Produkte $D_m \overline{D}_{n-m}$ zusammenfassen lassen, und weiter $n! - m!(n-m)!$ Glieder vom Typus

$$\pm a_{1k_1} a_{2k_2} \dots a_{mk_m} a_{m+1, k_{m+1}} \dots a_{nk_n},$$

wobei $k_1, k_2, \dots, k_m$ nicht, abgesehen von der Reihenfolge, mit $1, 2, \dots, m$ übereinstimmen und infolgedessen auch noch unter den Indizes $k_{m+1}, \dots, k_n$ jedesmal wenigstens eine der Zahlen $1, 2, \dots, m$ sich findet. Nach den Beziehungen $a_{hk} = a_{kh}$ und den Ungleichungen (13) erweist sich nun ein jedes solches Glied dem Betrage nach

$$\leq \frac{1}{4} a_{11} a_{22} \dots a_{mm} a_{mm} a_{m+2, m+2} \dots a_{nn}.$$

Benutzen wir die Ungleichungen (16) und (17) und verstehen unter λ_n eine in der Folge noch zu fixierende positive Konstante, so entsteht jetzt

$$D - \lambda_n a_{11} a_{22} \dots a_{nn} \geq \{(\lambda_m \lambda_{n-m} - \lambda_n) a_{m+1, m+1} - \frac{1}{4} (n! - m!(n-m)!) a_{mm}\} a_{11} a_{22} \dots a_{mm} a_{m+2, m+2} \dots a_{nn}.$$

Wir können

$$1 = \lambda_1 > \lambda_2 > \dots > \lambda_{n-1}$$

voraussetzen. Die positive Größe λ_n denken wir uns zunächst irgendwie, jedoch kleiner als jedes der Produkte $\lambda_m \lambda_{n-m}$ für $m = 1, 2, \dots, n-1$ gewählt und setzen zur Abkürzung

$$\frac{1}{4} \frac{(n! - m!(n-m)!)}{\lambda_m \lambda_{n-m} - \lambda_n} = \kappa_{m+1} \quad (m = 1, 2, \dots, n-1).$$

Alsdann wird die Ungleichung (15) mit dem betreffenden Werte λ_n jedenfalls schon statthaben, wenn für wenigstens einen der Indizes $m = 1, 2, \dots, n-1$ sich

$$a_{m+1, m+1} \geq \kappa_{m+1} a_{m m}$$

herausstellt.

6. Nunmehr haben wir uns nur noch mit der *Annahme* zu beschäftigen, daß sämtliche Ungleichungen

$$(18) \quad \kappa_2 a_{11} > a_{22}, \kappa_3 a_{22} > a_{33}, \dots, \kappa_n a_{n-1, n-1} > a_{nn}$$

gelten.

Aus den Ungleichungen (16) geht unter Beachtung der Ungleichungen (14) und von $a_{11} > 0$ hervor, daß sicherlich die *Determinanten*

$$D_1, D_2, \dots, D_{n-1}$$

positiv ausfallen und daher auch

$$q_1 = D_1, q_2 = \frac{D_2}{D_1}, \dots, q_{n-1} = \frac{D_{n-1}}{D_{n-2}}$$

sämtlich endlich und > 0 sind. Wir können daher, mag nun die Determinante $D_n > 0$ und damit auch f positiv sein oder nicht, jedenfalls bereits für f die Darstellung aus § 1:

$$(19) \quad f = q_1 \xi_1^2 + q_2 \xi_2^2 + \dots + q_n \xi_n^2,$$

$$(20) \quad \xi_1 = x_1 + \gamma_{12} x_2 + \dots + \gamma_{1n} x_n, \xi_2 = x_2 + \dots + \gamma_{2n} x_n, \dots, \xi_n = x_n$$

mit

$$q_n = \frac{D_n}{D_{n-1}}, \quad \gamma_{hk} = \frac{D \left(\begin{smallmatrix} 1, \dots, h-1, h \\ 1, \dots, h-1, k \end{smallmatrix} \right)}{D_h} \quad \left(\begin{matrix} h = 1, \dots, n-1 \\ k = h+1, \dots, n \end{matrix} \right)$$

ansetzen.

Die Determinante, welche den Zähler des vorstehenden Quotienten für γ_{hk} bildet, besteht aus $h!$ Gliedern, von denen ein jedes mit Rücksicht auf die Ungleichungen (13) sich dem Betrage nach $\leq \frac{1}{2} a_{11} a_{22} \dots a_{hh}$ erweist, während der Nenner dieses Quotienten nach (16) sich $\geq \lambda_h a_{11} a_{22} \dots a_{hh}$ findet. Danach hat man

$$(21) \quad |\gamma_{hk}| \leq \frac{1}{2} \frac{h!}{\lambda_h} \quad \left(\begin{matrix} h = 1, \dots, n-1 \\ k = h+1, \dots, n \end{matrix} \right).$$

Aus (19) und (20) geht

$$q_1 = a_{11}, q_2 \leq a_{22}, \dots, q_{n-1} \leq a_{n-1, n-1}$$

hervor, woraus mit Rücksicht auf (18)

$$(22) \quad q_1 = a_{11}, q_2 < \kappa_2 a_{11}, \dots, q_{n-1} < \kappa_2 \kappa_3 \dots \kappa_{n-1} a_{11}$$

folgt.

7. Es kommt jetzt vor allem darauf an, zu erkennen, daß wir von den Ungleichungen (I) eine endliche Anzahl herausgreifen können, infolge deren sich auch q_n als > 0 und damit also f als eine wesentlich positive Form erweist. Für diesen wichtigen Nachweis können wir uns des elementaren Prinzips bedienen, welches Dirichlet so meisterhaft in der Theorie der algebraischen Einheiten gehandhabt hat, wonach bei Verteilung einer Anzahl von Größensystemen in eine kleinere Anzahl von Bereichen darunter irgendein Bereich da sein muß, der wenigstens zwei der Systeme auf einmal aufnimmt.

Wir bilden zum Ausdrucke (19) von f mit den nämlichen $\xi_1, \xi_2, \dots, \xi_n$ die Hilfsform

$$(23) \quad F = \xi_1^2 + \kappa_2 \xi_2^2 + \dots + \kappa_2 \kappa_3 \dots \kappa_{n-1} \xi_{n-1}^2 + \mu \xi_n^2,$$

wobei μ folgende Bedeutung haben soll. Es seien $t_1, t_2, \dots, t_{n-1}$ beziehlich die kleinsten ganzen Zahlen, für welche

$$(24) \quad t_1^2 \geq n, t_2^2 \geq n\kappa_2, \dots, t_{n-1}^2 \geq n\kappa_2 \kappa_3 \dots \kappa_{n-1}$$

ausfällt, und es werde

$$(25) \quad \mu = \frac{1}{n t_1^2 t_2^2 \dots t_{n-1}^2}$$

gesetzt.

Danach ist μ eine bestimmte positive Größe, mithin F eine positive Form der Variablen $x_1, x_2, \dots, x_n$.

Wir zeigen, daß man für $x_1, x_2, \dots, x_n$ ganze Zahlen, unter denen $x_n \neq 0$ ist, finden kann, so daß dafür $F \leq 1$ ausfällt.

In der Tat, zunächst ist $\xi_n = x_n$. Wir setzen der Reihe nach $x_n = 0, 1, 2, \dots, t_1 t_2 \dots t_{n-1}$ und können zu jedem dieser Werte von x_n nacheinander $x_{n-1}, x_{n-2}, \dots, x_1$ als ganze Zahlen derart bestimmen, daß

$$(26) \quad 0 \leq \xi_{n-1} < 1, 0 \leq \xi_{n-2} < 1, \dots, 0 \leq \xi_1 < 1$$

wird. Wir erhalten damit $t_1 t_2 \dots t_{n-1} + 1$ in x_n durchweg verschiedene Wertsysteme $x_1, x_2, \dots, x_n$, welche sämtlich diese Bedingungen (26) erfüllen.

Zerlegen wir nun für $h = n - 1, n - 2, \dots, 1$ jedesmal das Intervall $0 \leq \xi_h < 1$ in die t_h Intervalle

$$0 \leq \xi_h < \frac{1}{t_h}, \frac{1}{t_h} \leq \xi_h < \frac{2}{t_h}, \dots, \frac{t_h - 1}{t_h} \leq \xi_h < 1,$$

so tritt eine Teilung des ganzen durch (26) definierten Gebietes in $t_1 t_2 \dots t_{n-1}$ völlig untereinander getrennte Gebiete ein und wird man unter

jenen Systemen $x_1, x_2, \dots, x_n$, deren Anzahl $> t_1 t_2 \dots t_{n-1}$ ist, gewiß irgend zwei, etwa

$$x'_1, x'_2, \dots, x'_n; \quad x''_1, x''_2, \dots, x''_n \quad (x''_n > x'_n)$$

finden können, für welche $\xi_1, \xi_2, \dots, \xi_{n-1}$ demselben Teilgebiete angehören.

Für das durch Subtraktion aus beiden hervorgehende System

$$x_1 = x''_1 - x'_1, \quad x_2 = x''_2 - x'_2, \quad \dots, \quad x_n = x''_n - x'_n$$

gelten dann die Ungleichungen

$$(27) \quad |\xi_1| < \frac{1}{t_1}, \dots, |\xi_{n-1}| < \frac{1}{t_{n-1}}, \quad |\xi_n| \leq t_1 t_2 \dots t_{n-1},$$

während zugleich $x_n > 0$ ist. Aus (27), (24) und (25) ersieht man, daß infolgedessen weiter die $n-1$ ersten Terme des Ausdrucks F ein jeder $< \frac{1}{n}$, der $n^{\text{te}} \leq \frac{1}{n}$ werden und also für das betreffende ganzzahlige Wertsystem $x_1, x_2, \dots, x_n$ mit positivem x_n der ganze Ausdruck F gewiß < 1 ausfällt.

Das erhaltene System $x_1, x_2, \dots, x_n$ befreien wir noch von einem in den Zahlen etwa vorhandenen gemeinsamen Teiler > 1 , wobei die Ungleichung $F \leq 1$ nicht verloren geht.

Andererseits können wir, *allein in Berücksichtigung* der Ausdrücke (20) für $\xi_1, \xi_2, \dots, \xi_n$ und der Ungleichungen (21) für die Koeffizienten γ_{hk} , von vornherein eine endliche Anzahl bestimmter ganzzahliger Systeme $x_1, x_2, \dots, x_n$ mit positivem x_n und ohne gemeinsamen Teiler > 1 anweisen, welche überhaupt als die einzigen derartigen Systeme in Frage kommen, die vielleicht den Ungleichungen (27) genügen können. Wir ermitteln die sämtlichen bezüglichen Systeme und wir heben nunmehr *viertens* unter den Ungleichungen (I) diejenigen heraus, welche ausdrücken, daß für diese besonderen Systeme stets $f(x_1, x_2, \dots, x_n) \geq a_{11}$ sein soll.

Als dann muß infolge aller bereits herausgehobenen Ungleichungen (I) notwendig

$$(28) \quad q_n \geq \mu a_{11}$$

werden. Denn in dem Ausdrücke (19) von f sind wegen (22) die Koeffizienten von $\xi_1^2, \xi_2^2, \dots, \xi_{n-1}^2$ durchweg nicht größer als in dem Multiplum $a_{11}F$ von F . Wäre nun $q_n < \mu a_{11}$, so würde daher bei von Null verschiedenem x_n stets $f < a_{11}F$ sein. Wir haben aber ein solches ganzzahliges System $x_1, x_2, \dots, x_n$ mit von Null verschiedenem x_n , wofür $F \leq 1$ ist, während wir für das betreffende System hier bereits einer der Ungleichungen (I) gemäß $a_{11} \leq f$ fordern, also ergäbe sich ein Widerspruch und folgt in der Tat die Ungleichung (28).

8. Indem wir (28) mit allen Ungleichungen (18) multiplizieren, geht

$$q_n > \frac{\mu}{x_2 x_3 \dots x_n} a_{nn}$$

hervor, und indem wir weiter diese Ungleichung mit der Ungleichung für D_{n-1} nach (16) multiplizieren, folgt

$$D_n > \frac{\lambda_{n-1} \mu}{\kappa_2 \kappa_3 \dots \kappa_n} a_{11} a_{22} \dots a_{nn}.$$

Nach der Bedeutung von μ , die aus (24) und (25) ersichtlich ist, können wir nun über λ_n *definitiv* in solcher Art verfügen, daß *hieraus wieder die Ungleichung* (15) *für* D_n *hervorgeht*. Wir brauchen dazu nur λ_n kleiner als die $n-1$ Größen $\lambda_m \lambda_{n-m}$ ($m < n$) derart zu wählen, daß

$$\lambda_{n-1} \geq \lambda_n n \kappa_2 \kappa_3 \dots \kappa_n ([\sqrt{n}] + 1)^2 ([\sqrt{n \kappa_2}] + 1)^2 \dots ([\sqrt{n \kappa_2 \dots \kappa_{n-1}}] + 1)^2$$

ist. Die Klammer $[\]$ soll hier das bekannte Zeichen für größte Ganze vorstellen. Dieser Forderung für λ_n aber können wir immer entsprechen, weil der Ausdruck rechts hier nach der Art, wie $\kappa_2, \dots, \kappa_n$ von λ_n abhängen, gleichzeitig mit λ_n abnimmt und der Null zustrebt.

Bei dieser Bestimmungsweise von λ_n gilt nunmehr infolge der bisher herausgehobenen Ungleichungen (I) *in allen Fällen die Ungleichung* (15) *für* D_n .

9. Indem wir auf den Zähler einer Größe $q_h = D_h : D_{h-1}$ die Ungleichung (16) bzw. (15), auf den Nenner die Ungleichung

$$D_{h-1} \leq a_{11} a_{22} \dots a_{h-1, h-1}$$

in Anwendung bringen, entsteht allgemein

$$q_h \geq \lambda_h a_{hh} \quad (h = 1, 2, \dots, n),$$

und nun zeigt sich in der Tat leicht, daß eine endliche Anzahl unter den Ungleichungen (I) angewiesen werden kann, welche alle übrigen dieser Ungleichungen zur Folge haben.

Wir setzen für jeden Wert $l = 1, 2, \dots, n$ die folgenden Ungleichungen an:

$$(29) \quad \lambda_n \xi_n^2 < 1, \lambda_{n-1} \xi_{n-1}^2 < 1, \dots, \lambda_l \xi_l^2 < 1,$$

$$(30) \quad \lambda_{l-1} \xi_{l-1}^2 < \frac{l-1}{4}, \lambda_{l-2} \xi_{l-2}^2 < \frac{l-2}{4}, \dots, \lambda_2 \xi_2^2 < \frac{2}{4}, \lambda_1 \xi_1^2 \leq \frac{1}{4}$$

mit den Ausdrücken

$$(31) \quad \xi_1 = x_1 + \gamma_{12} x_2 + \dots + \gamma_{1n} x_n, \xi_2 = x_2 + \dots + \gamma_{2n} x_n, \dots, \xi_n = x_n$$

und den Einschränkungen

$$(32) \quad |\gamma_{hk}| \leq \frac{1}{2} \frac{h!}{\lambda_h}.$$

Es ist einleuchtend, daß diesen Bedingungen jedesmal nur eine *endliche* Anzahl ganzzahliger Systeme $x_1, x_2, \dots, x_n$ entsprechen können. *Unter diesen wählen wir alle Systeme* $s_1^{(l)}, s_2^{(l)}, \dots, s_n^{(l)}$ *aus, in welchen* $s_1^{(l)}, s_{l+1}^{(l)}, \dots, s_n^{(l)}$ *keinen gemeinsamen Teiler* > 1 *haben, und fordern jedesmal die Bedingung* (I) *für die betreffenden Systeme.*

Die in solcher Weise bisher hervorgehobenen der Ungleichungen (I) ziehen dann notwendig die Gesamtheit dieser unendlich vielen Ungleichungen nach sich.

In der Tat, es sei $x_1 = s_1^{(l)}, \dots, x_n = s_n^{(l)}$ ein ganzzahliges System, welches nicht zu den eben hervorgehobenen gehört und wobei $s_1^{(l)}, \dots, s_n^{(l)}$ keinen gemeinsamen Teiler > 1 haben. Alsdann bestehen dafür insbesondere mit den zu f gehörigen Ausdrücken $\xi_1, \xi_2, \dots, \xi_n$ nicht alle Ungleichungen (29), (30).

Wir nehmen *erstens* an, es sei dafür etwa $\lambda_h \xi_h^2 \geq 1$ und $h \geq l$. Dann folgt aus

$$f = q_1 \xi_1^2 + q_2 \xi_2^2 + \dots + q_n \xi_n^2$$

mit Rücksicht auf $q_h \geq \lambda_h a_{hh}$ und $a_{hh} \geq a_{ll}$ sofort $f \geq a_{ll}$.

Es seien *zweitens* für das bezügliche System in f sämtliche Ungleichungen (29) erfüllt, und es sei $h (< l)$ der *größte* Index, für den statt der in (30) genannten Ungleichung vielmehr sich die Ungleichung $\lambda_h \xi_h^2 \geq \frac{1}{4} h$ bzw. $\xi_1^2 > \frac{1}{4}$, je nachdem $h > 1$ oder $= 1$ ist, einstellt. Dann haben wir wegen $q_h \geq \lambda_h a_{hh}$ für das betreffende System $x_1, \dots, x_n$ jedenfalls

$$(33) \quad f(x_1, \dots, x_n) \geq \frac{1}{4} (a_{11} + a_{22} + \dots + a_{hh}) + q_{h+1} \xi_{h+1}^2 + \dots + q_n \xi_n^2.$$

Wir denken uns nun die Zahlen $x_{h+1}, \dots, x_n$ festgehalten, statt der Zahlen $x_h, x_{h-1}, \dots, x_1$ aber, was offenbar möglich ist, nacheinander solche ganze Zahlen $x_h^*, x_{h-1}^*, \dots, x_1^*$ gewählt, daß die neuen dazugehörigen Ausdrücke (31) die Bedingungen

$$(34) \quad -\frac{1}{2} \leq \xi_h^* \leq \frac{1}{2}, -\frac{1}{2} \leq \xi_{h-1}^* \leq \frac{1}{2}, \dots, -\frac{1}{2} \leq \xi_1^* \leq \frac{1}{2}$$

erfüllen. Wir haben damit ein modifiziertes ganzzahliges System $x_1^*, \dots, x_n^*$ erhalten, in welchem $x_i^* = x_i, \dots, x_n^* = x_n$ keinen gemeinsamen Teiler > 1 haben, und welches nunmehr *allen* Bedingungen (29), (30) genügt, für das wir also bereits $f \geq a_{ll}$ voraussetzen. Jetzt folgt aus (33) und (34), da stets $a_{kk} \geq q_k$ ist,

$$f(x_1, \dots, x_n) \geq f(x_1^*, \dots, x_n^*) \geq a_{ll}.$$

10. Wir können unser Endergebnis folgendermaßen aussprechen:

Der reduzierte Raum B ist ein konvexer Kegel mit der Spitze im Nullpunkte $f = 0$, der von einer endlichen Anzahl durch diesen Punkt laufender Ebenen begrenzt wird.

§ 6. Die Kanten des reduzierten Raumes.

1. Im ganzen reduzierten Raume B gilt $a_{11} \geq 0$; die Punkte darin, für welche $a_{11} > 0$ ist, entsprechen positiven Formen; die Punkte darin,

für welche $a_{11} = 0$ ist, erfüllen zugleich die Bedingungen $a_{12} = 0$, $a_{13} = 0$, $\dots$, $a_{1n} = 0$ und bilden eine nur $\frac{n(n-1)}{2}$ -fache Mannigfaltigkeit auf der Begrenzung von B .

Die zur Charakterisierung des reduzierten Raumes dienenden Ungleichungen haben eine jede die Gestalt

$$(I, II) \quad \sum m_{hk} a_{hk} \geq 0,$$

worin die m_{hk} gewisse gegebene numerische Koeffizienten vorstellen. Daraus geht hervor:

Ist f eine reduzierte Form, so ist auch jedes Produkt cf , wo c einen positiven Faktor vorstellt, sind f und g zwei reduzierte Formen, so ist auch jede Verbindung $(1-t)f + tg$, wo $0 < t < 1$ ist, eine reduzierte Form.

2. Die allgemeinen Prinzipien über lineare Ungleichungen, welche ich in meiner „Geometrie der Zahlen“ § 19 dargelegt habe, ergeben folgende Aufschlüsse über die zur Definition des Raumes B wirklich notwendigen von den sämtlichen Ungleichungen (I) und (II).

Wir wollen eine nicht identisch verschwindende reduzierte Form φ eine Kantenform des reduzierten Raumes nennen, wenn es nicht möglich ist, φ als Summe zweier reduzierter Formen darzustellen, die weder identisch verschwinden, noch positive Vielfache voneinander sind.

Eine Kantenform ist vollständig zu charakterisieren als eine reduzierte Form, für welche unter den Ungleichungen (I, II) irgend $\frac{n(n+1)}{2} - 1$ solche, deren linke Seiten linear unabhängige Funktionen der a_{hk} sind, mit dem Zeichen $=$ erfüllt sind.

Gelten uns die Formen, die Vielfache voneinander sind, als nicht wesentlich verschieden, so gibt es nur eine endliche Anzahl wesentlich verschiedener Kantenformen. Die Strahlen vom Nullpunkte nach ihnen bilden die Kanten des reduzierten Raumes.

Von den Ungleichungen (I, II) ist zur Definition des reduzierten Raumes nur jede solche wesentlich, die mit dem Zeichen $=$ eine Ebene liefert, welche $\frac{n(n+1)}{2} - 1$ unabhängige Kanten des reduzierten Raumes aufnimmt, d. h. solche $\frac{n(n+1)}{2} - 1$ Kanten, durch die sich nur eine einzige Ebene legen läßt. Die so charakterisierten Ungleichungen bestimmen die Wände des reduzierten Raumes. Alle übrigen Ungleichungen (I, II) können bei der Definition des reduzierten Raumes als aus diesen Ungleichungen folgend fortgelassen werden.

3. Wir greifen auf jeder Kante des reduzierten Raumes eine Form heraus, etwa diejenige, für welche der erste von Null verschiedene unter den

Koeffizienten $a_{11}, a_{22}, \dots, a_{nn}$ den Wert 1 hat. Wir erhalten auf diese Weise eine endliche Anzahl völlig bestimmter Formen $\varphi_1, \varphi_2, \dots, \varphi_r$. Indem der reduzierte Raum ein konvexer Kegel vom Nullpunkte aus ist, besteht alsdann der Satz:

Jede reduzierte Form f läßt sich (auf eine oder auf unendlich viele Weisen) in die Gestalt

$$(35) \quad f = c_1 \varphi_1 + c_2 \varphi_2 + \dots + c_r \varphi_r$$

mit Koeffizienten $c_1, c_2, \dots, c_r$, die sämtlich ≥ 0 sind, setzen, und umgekehrt ist jede Form, die sich in dieser Weise darstellen läßt, eine reduzierte.

§ 7. Die Nachbarkammern des reduzierten Raumes.

Wir beweisen in betreff der reduzierten Formen noch den Satz:

Die ganzzahligen Substitutionen von der Determinante ± 1 , welche fähig sind, positive reduzierte Formen wieder in reduzierte Formen überzuführen, existieren nur in endlicher Anzahl.

Wir können diesen Satz auch folgendermaßen aussprechen:

Im Gebiete der positiven Formen grenzt der reduzierte Raum nur an eine endliche Anzahl von den äquivalenten Kammern an.

1. Es sei S :

$$x_h = s_{h1}y_1 + s_{h2}y_2 + \dots + s_{hn}y_n \quad (h = 1, 2, \dots, n)$$

eine unimodulare ganzzahlige Substitution, welche

$$f = \sum a_{hk}x_hx_k \text{ in } g = \sum b_{hk}y_hy_k$$

überführt, wobei beide Formen reduziert sind und $a_{11} > 0$ ist.

Wir haben dabei

$$(36) \quad f(s_{1k}, s_{2k}, \dots, s_{nk}) = b_{kk} \quad (k = 1, 2, \dots, n).$$

Ist die letzte der Zahlen $s_{1k}, s_{2k}, \dots, s_{nk}$, die von Null verschieden ist, s_{hk} , so ergibt sich, da ihr absoluter Betrag mindestens 1 sein muß, bei Gebrauch der früheren Bezeichnungen q_h in bezug auf f :

$$b_{kk} \geq q_h \geq \lambda_h a_{hh} \geq \lambda_n a_{hh}.$$

2. Daraus schließen wir zunächst, daß für jeden Index l durchaus

$$b_{ll} \geq \lambda_n a_{ll} \quad (l = 1, 2, \dots, n)$$

sein muß.*)

Denn hätte man für einen Index l :

$$b_{ll} < \lambda_n a_{ll},$$

so würde daraus weiter

$$b_{11} \leq b_{22} \leq \dots \leq b_{ll} < \lambda_n a_{ll} \leq \lambda_n a_{l+1, l+1} \leq \dots \leq \lambda_n a_{nn}$$

*) Eine ähnliche Überlegung findet sich bei C. Jordan, Journal de l'École polytechnique, T. XXIX, cahier 48, p. 127.

hervorgehen und daher könnte keine Größe

$$s_{hk} (h = l, l+1, \dots, n; k = 1, 2, \dots, l)$$

in ihrer, der l^{ten} Vertikalreihe die letzte von Null verschiedene Zahl sein, d. h. diese Größen müßten *sämtlich Null* sein. Sie bilden aber in der Determinante der Substitution S die Elemente, die gleichzeitig in bestimmten l Vertikal- und $n-l+1$ Horizontalreihen vorkommen, also wäre diese Determinante Null, während sie ± 1 sein sollte.

Ganz entsprechend wird, weil auch g durch eine unimodulare ganzzahlige Substitution in f übergeht, *stets*

$$(37) \quad a_{kk} \geq \lambda_n b_{kk} \quad (k = 1, 2, \dots, n)$$

sein.

3. Wir nehmen nun erstens an, für jeden Index $k = 1, 2, \dots, n-1$ fände sich *stets*

$$(38) \quad b_{kk} \geq \lambda_n a_{k+1, k+1},$$

so folgt daraus vermöge (37) allgemein

$$a_{kk} \geq \lambda_n^2 a_{k+1, k+1}$$

und weiter

$$a_{11} \geq \lambda_n^{2k-2} a_{kk} \geq \lambda_n^{2k-1} b_{kk} \quad (k = 1, 2, \dots, n).$$

Die zu den Zahlen $x_1 = s_{1k}, \dots, x_n = s_{nk}$ gehörenden Werte von $\xi_1, \dots, \xi_n$ erfüllen dann nach (36) und da allgemein $q_h \geq \lambda_n a_{hh} \geq \lambda_n a_{11}$ ist, die Ungleichung

$$\lambda_n^{2k} (\xi_1^2 + \xi_2^2 + \dots + \xi_n^2) \leq 1,$$

und hieraus ist zu ersehen, daß für jede Vertikalreihe $s_{1k}, \dots, s_{nk}$ von S nur eine *endliche* Anzahl ganzzahliger Systeme in Betracht kommen.

4. Ist die Annahme (38) nicht immer zutreffend, so sei l der *größte Index, wofür*

$$(39) \quad b_{l-1, l-1} < \lambda_n a_{11} \quad (l \geq 2)$$

ist. Wir haben dann nacheinander *viererlei* Umstände in Betracht zu ziehen.

Erstens zeigt eine ähnliche Überlegung wie in 2., daß jedenfalls *alle Koeffizienten*

$$s_{hk} (h = l, l+1, \dots, n; k = 1, 2, \dots, l-1)$$

Null sind. Die Substitution S kann also geschrieben werden:

$$x_h = s_{h1}y_1 + \dots + s_{h, l-1}y_{l-1} + s_{hl}y_l + \dots + s_{hn}y_n \quad (h < l),$$

$$x_h = \quad \quad \quad s_{hl}y_l + \dots + s_{hn}y_n \quad (h \geq l).$$

Jetzt ist

$$x_h = s_{h1}y_1 + \dots + s_{h, l-1}y_{l-1} \quad (h = 1, 2, \dots, l-1)$$

eine unimodulare ganzzahlige Substitution von $l-1$ Variablen, und durch sie geht die positive Form

$$f(x_1, \dots, x_{l-1}, 0, \dots, 0) \quad \text{in} \quad g(y_1, \dots, y_{l-1}, 0, \dots, 0)$$

über, wobei beide Formen reduzierte von $l - 1$ Variablen sind. Nehmen wir den zu beweisenden Satz, der für Formen von einer Variable evident ist, als bereits bewiesen für Formen mit weniger als n Variablen an, so kommen danach *zweitens für die Koeffizienten*

$$s_{hk} (h = 1, 2, \dots, l - 1; k = 1, 2, \dots, l - 1)$$

nur eine endliche Anzahl von Systemen in Betracht.

Unsere Annahme über die Zahl l schließt in sich, daß, falls $l < n$ ist, die Beziehungen

$$b_{kk} \geq \lambda_n a_{k+1, k+1} \quad (k = l, l + 1, \dots, n - 1)$$

gelten, woraus wir vermöge (37) weiter

$$a_{kk} \geq \lambda_n^2 a_{k+1, k+1} \quad (k = l, l + 1, \dots, n - 1)$$

und sodann

$$a_{ii} \geq \lambda_n^{2k-2i} a_{kk} \geq \lambda_n^{2k-2i+1} b_{kk} \quad (k = l, l + 1, \dots, n - 1, n)$$

entnehmen; letztere Beziehung besteht auch für $l = n$, $k = n$.

Die Gleichung (36) für b_{kk} liefert daraufhin für die zu $s_{ik}, s_{i+1, k}, \dots, s_{nk}$ gehörenden Werte $\xi_i, \xi_{i+1}, \dots, \xi_n$ die Relation

$$\lambda_n^{2k-2i+2} (\xi_i^2 + \xi_{i+1}^2 + \dots + \xi_n^2) \leq 1.$$

Hieraus ist *drittens* zu ersehen, daß *für die Zahlen*

$$s_{hk} (h = l, l + 1, \dots, n; k = l, l + 1, \dots, n)$$

nur eine endliche Anzahl von Systemen in Betracht kommen.

Endlich muß g als reduzierte Form insbesondere allen Ungleichungen

$$g(y_1, \dots, y_{l-1}, e_i^{(k)}, \dots, e_n^{(k)}) \geq b_{kk} \quad (k = l, l + 1, \dots, n)$$

genügen, in welchen $e_k^{(k)} = 1$, die anderen Zahlen $e_h^{(k)} = 0$ und $y_1, \dots, y_{l-1}$ beliebige ganze Zahlen sind. Diese Ungleichungen kommen nach der bereits festgestellten besonderen Gestalt der Substitution S auf die sämtlichen Ungleichungen

$$f(x_1, \dots, x_{l-1}, s_{ik}, \dots, s_{nk}) \geq f(s_{1k}, \dots, s_{l-1, k}, s_{ik}, \dots, s_{nk}) \quad (k = l, l + 1, \dots, n)$$

für beliebige ganze Zahlen $x_1, \dots, x_{l-1}$ hinaus.

Eine ähnliche Überlegung, wie sie am Schlusse von § 5 zur Anwendung kam, ergibt nun, daß hiernach die zu $s_{1k}, \dots, s_{nk} (k \geq l)$ gehörenden Verbindungen $\xi_{l-1}, \dots, \xi_l$ jedenfalls die Bedingungen

$$\frac{1}{4} (q_1 + \dots + q_h) \geq q_1 \xi_1^2 + \dots + q_h \xi_h^2 \quad (h = l - 1, \dots, 1)$$

und umsomehr die Bedingungen

$$\lambda_n \xi_{l-1}^2 \leq \frac{l-1}{4}, \dots, \lambda_n \xi_1^2 \leq \frac{1}{4}$$

erfüllen, und hiernach kommen *viertens* auch *für die Zahlen*

$$s_{hk} (h = 1, 2, \dots, l - 1; k = l, l + 1, \dots, n)$$

nur eine endliche Anzahl von Systemen in Betracht.

Damit ist der zu führende Nachweis in allen Punkten erbracht.

§ 8. Die Determinantenfläche.

Der Raum der positiven Formen wird von der *Flächenschar* $D(f) = \text{const.}$ durchzogen, deren einzelne Flächen jedesmal alle Formen f zu einem und dem nämlichen positiven Determinantenwert aufnehmen. Alle Flächen dieser Schar sind untereinander ähnlich und ähnlich gelegen vom Nullpunkte aus, so daß es für die meisten Zwecke genügt, von ihnen etwa die eine Fläche $D(f) = 1$ in Betracht zu ziehen.

Wir beweisen hier folgenden Satz:

Sind f und g zwei verschiedene positive Formen von der Determinante 1 und ist t ein beliebiger Wert > 0 und < 1 , so hat die gleichfalls positive Form $(1-t)f + tg$ stets eine Determinante > 1 .

Wir können bekanntlich (sowie von den Formen f und g auch nur eine positiv ist) immer eine *lineare Transformation* von der Determinante 1 mit *lauter reellen Koeffizienten* finden, wodurch f und g gleichzeitig in *Aggregate*

$$\alpha_1 z_1^2 + \alpha_2 z_2^2 + \dots + \alpha_n z_n^2, \quad \beta_1 z_1^2 + \beta_2 z_2^2 + \dots + \beta_n z_n^2$$

übergehen, welche nur die Quadrate der neuen Variablen enthalten. Die Determinante der Verbindung $(1-t)f + tg$ gewinnt dann allgemein den Produktausdruck

$$\Delta(t) = (\alpha_1 + t(\beta_1 - \alpha_1))(\alpha_2 + t(\beta_2 - \alpha_2)) \dots (\alpha_n + t(\beta_n - \alpha_n)),$$

und nach Voraussetzung ist

$$\Delta(0) = \alpha_1 \alpha_2 \dots \alpha_n = 1, \quad \Delta(1) = \beta_1 \beta_2 \dots \beta_n = 1.$$

Nun erhalten wir

$$\frac{d^2 \log \Delta(t)}{dt^2} = - \left(\frac{\beta_1 - \alpha_1}{\alpha_1 + t(\beta_1 - \alpha_1)} \right)^2 - \dots - \left(\frac{\beta_n - \alpha_n}{\alpha_n + t(\beta_n - \alpha_n)} \right)^2.$$

Dieser Ausdruck fällt danach im ganzen Intervalle $0 \leq t \leq 1$ stets < 0 aus, d. h. die Kurve $u = \log \Delta(t)$ in einer t, u -Ebene ist im Bereich $0 \leq t \leq 1$ stets konvex von der t -Achse fort. Mithin ist diese Funktion u , da sie an den zwei Endpunkten des Intervalls $= 0$ ist, im Inneren desselben durchweg > 0 , also ist hier stets $\Delta(t) > 1$, was zu zeigen war.

Setzen wir

$$c(1-t)\alpha_h = a_h, \quad ct\beta_h = b_h \quad (h = 1, 2, \dots, n),$$

wo c ein positiver Faktor sei, so kommt die in betreff des Ausdrucks $\Delta(t)$ bewiesene Ungleichung auf den Satz hinaus:

Sind $a_1, a_2, \dots, a_n$ und $b_1, b_2, \dots, b_n$ lauter positive Größen, ohne daß man gerade $a_1 : a_2 : \dots : a_n = b_1 : b_2 : \dots : b_n$ hat, so gilt stets die Beziehung

$$\sqrt[n]{(a_1 + b_1)(a_2 + b_2) \dots (a_n + b_n)} > \sqrt[n]{a_1 a_2 \dots a_n} + \sqrt[n]{b_1 b_2 \dots b_n}.$$

Das über die Fläche $D(f) = 1$ gefundene Resultat können wir auch folgendermaßen ausdrücken:

Konstruiert man in irgendeinem Punkte der Determinantenfläche $D(f) = 1$, welcher einer positiven Form f entspricht, die Tangentialebene an diese Fläche, so liegt im ganzen Bereiche der positiven Formen diese Fläche, abgesehen vom Berührungspunkte, vollständig auf der dem Nullpunkte abgewandten Seite der Ebene.

Indem wir diese Lage der Tangentialebene an $D(f) = \text{const.}$ durch einen Punkt f in bezug auf irgendeinen zweiten Punkt g der Fläche in eine Formel fassen, kommen wir auf Grund der schon oben verwandten gleichzeitigen Transformation von f und g in Aggregate von Quadraten zu

$$(40) \quad \frac{1}{n} \left(\frac{\beta_1}{\alpha_1} + \frac{\beta_2}{\alpha_2} + \dots + \frac{\beta_n}{\alpha_n} \right) > \sqrt[n]{\frac{\beta_1 \beta_2 \dots \beta_n}{\alpha_1 \alpha_2 \dots \alpha_n}},$$

d. i. einfach zu der bekannten Ungleichung zwischen dem arithmetischen und geometrischen Mittel von n positiven, nicht lauter gleichen Größen.

Kürzer können wir den Charakter der Fläche $D(f) = 1$ noch dahin schildern:

Die Determinantenfläche $D(f) = 1$ ist im Gebiete der positiven Formen überall konvex nach dem Nullpunkte zu.

§ 9. Das Problem der dichtesten gitterförmigen Lagerung von Kugeln.

1. Für den Koeffizienten a_{11} einer positiven reduzierten Form f haben wir stets

$$f(x_1, x_2, \dots, x_n) \geq a_{11},$$

wenn $x_1, x_2, \dots, x_n$ ganze Zahlen ohne gemeinsamen Teiler > 1 sind, und diese Ungleichung überträgt sich sofort auf beliebige Systeme von ganzen Zahlen $x_1, x_2, \dots, x_n$, die nur nicht sämtlich Null sind. Danach ist a_{11} die kleinste durch die Form f mittels ganzer, nicht sämtlich verschwindender Zahlen darstellbare Größe. Wir nennen diese Größe das Minimum der Form f und schreiben sie $M(f)$; sie ist (wie die ganze reduzierte Form) offenbar eine Invariante der Klasse f .

2. Aus den Ungleichungen (14) und (15) in § 5 (nach der Größenfolge von $a_{11}, a_{22}, \dots, a_{nn}$ und der oberen Begrenzung ihres Produktes mit Rücksicht auf die Determinante) entnehmen wir

$$(41) \quad D(f) \geq \lambda_n (M(f))^n.$$

Danach überschreitet $\frac{M(f)}{\sqrt[n]{D(f)}}$ niemals eine gewisse nur von n abhängende Grenze.

Es ist nun durch Hermite die Frage nach dem präzisen Maximum der Werte dieses Quotienten gestellt worden.

Korkine und Zolotareff*) definierten:

*) Mathematische Annalen, Bd. 6 und Bd. 11.

Eine positive quadratische Form f mit n Variablen soll eine *extreme Form* (ihre Klasse eine *extreme Klasse*) heißen, wenn bei den infinitesimalen Variationen der Form der Quotient $M(f) : \sqrt[n]{D(f)}$ niemals zunimmt.

Da wir bei den Variationen durch bloße Multiplikation der variierten Form mit passendem Faktor den Wert des Minimums konstant erhalten können, so dürfen wir auch sagen:

Eine positive Form ist *extrem*, wenn bei keiner infinitesimalen Variation der Form, welche das Minimum ungeändert läßt, die Determinante abnehmen kann.

3. Den extremen Formenklassen kommt eine bemerkenswerte geometrische Bedeutung zu.

Bringen wir eine positive quadratische Form $f(x_1, x_2, \dots, x_n)$ auf die Gestalt

$$f = \xi_1^2 + \xi_2^2 + \dots + \xi_n^2,$$

so daß $\xi_1, \xi_2, \dots, \xi_n$ n reelle lineare Formen in $x_1, x_2, \dots, x_n$ sind, und deuten $\xi_1, \xi_2, \dots, \xi_n$ als rechtwinklige Koordinaten in einem Raume $\mathfrak{R}_n$ von n Dimensionen, so können wir $f \leq 1$ als eine n -dimensionale Kugel vom Radius 1 in diesem Raume bezeichnen. Ihr Volumen in $\xi_1, \xi_2, \dots, \xi_n$ ist

$$\gamma_n = \frac{\pi^{\frac{n}{2}}}{\Gamma(1 + \frac{n}{2})}.$$

Die Punkte $x_1 = m_1, x_2 = m_2, \dots, x_n = m_n$, welche ganzzahligen Werten $m_1, m_2, \dots, m_n$ entsprechen, bilden in dem Raume $\mathfrak{R}_n$ ein *parallelepipedisches Punktsystem* (Gitter) mit der Dichtigkeit $\frac{1}{\sqrt{D(f)}}$. Nämlich die einzelnen Parallelepipede

$$m_h - \frac{1}{2} \leq x_h \leq m_h + \frac{1}{2} \quad (h = 1, 2, \dots, n)$$

mit diesen Punkten als Mittelpunkten erfüllen in ihrer Gesamtheit den Raum $\mathfrak{R}_n$ lückenlos, sind untereinander kongruent und enthalten jedesmal je einen Gitterpunkt bei einem Volumen je $= \sqrt{D(f)}$.

Die n -dimensionalen Kugeln vom Radius $\frac{1}{2}\sqrt{M(f)}$ um diese einzelnen Gitterpunkte werden nun, nach der Bedeutung von $M(f)$, derart liegen, daß sie untereinander nur in Punkten der Begrenzung zusammenstoßen. Infolgedessen wird insbesondere

$$(42) \quad \gamma_n \left(\frac{1}{2} \sqrt{M(f)} \right)^n < \sqrt{D(f)}$$

gelten, eine Ungleichung, die den Charakter der Ungleichung (41) trägt, aber jener vorzuziehen sein wird, wie λ_n in § 5 festgesetzt wurde.

Halten wir nun die Bedeutung der Koordinaten $\xi_1, \xi_2, \dots, \xi_n$ in $\mathfrak{R}_n$ fest und variieren die Koeffizienten von f derart, daß $M(f)$ unverändert

bleibt, so bleiben diese Kugeln hier in ihrer Größe sowie in dem Charakter, nicht ineinander einzudringen, erhalten, es ändert sich nur das parallelepipedische Gitter ihrer Mittelpunkte.

Die Frage nach dem Maximum von $M(f) : \sqrt[n]{D(f)}$ oder den extremen Formenklassen ist danach, geometrisch gefaßt, gleichbedeutend mit der Frage nach den dichtesten gitterförmigen Lagerungen von lauter gleichen Kugeln im Raume von n Dimensionen.

§ 10. Bestimmung der extremen Formenklassen.

Eine Reihe interessanter Eigenschaften der extremen Klassen haben bereits Korkine und Zolotareff nachgewiesen. Auf Grund der hier entwickelten Reduktionsmethode der positiven Formen läßt sich die Theorie der extremen Formen in sehr befriedigender Weise zum Abschluß bringen.

1. Ich bezeichne eine reduzierte Form als eine *extreme Form in bezug auf den reduzierten Raum*, wenn bei allen denjenigen infinitesimalen Variationen der Form, wobei die Form im reduzierten Raume verbleibt, der Quotient $\frac{M(f)}{\sqrt[n]{D(f)}}$ niemals zunimmt.

2. Ist f eine positive Form und legen wir an die durch f laufende Determinantenfläche $D(f) = \text{const.}$ im Punkte f die Tangentialebene, so läßt nach dem Satze in § 8 diese Ebene die Fläche im Gebiete der positiven Formen (und vom Punkte f selbst abgesehen) ganz auf der dem Nullpunkte abgewandten Seite liegen. Also nimmt in dieser Tangentialebene beim Fortgang vom Punkte f aus die Determinante stets ab.

3. Wenn nun f reduziert ist, aber nicht gerade eine Kantenform des reduzierten Raumes vorstellt, so können wir f als Summe zweier positiven reduzierten Formen φ, ψ darstellen, die nicht Vielfache voneinander sind. Es seien dann φ^*, ψ^* diejenigen Vielfachen von φ, ψ , welche in die genannte Tangentialebene fallen, so liegt f innerhalb der geradlinigen Strecke von φ^* nach ψ^* . Nun wächst entweder beim Fortgang von f aus nach einer Seite dieser Strecke hin der Koeffizient a_{11} , oder er ist auf dieser ganzen Strecke konstant, also wäre es immer möglich, f so zu variieren, daß $D(f)$ abnimmt und gleichzeitig $M(f)$ nicht abnimmt, mithin $M(f) : \sqrt[n]{D(f)}$ wächst, und f wäre jedenfalls keine extreme Form in bezug auf den reduzierten Raum.

Wir erhalten demnach das Resultat:

Eine extreme Form in bezug auf den reduzierten Raum kann nur eine Kantenform in diesem Raume sein.

4. Jetzt sei $f = (a_{hk})$ eine positive Kantenform des reduzierten Raumes und extrem in bezug auf diesen Raum, und wir legen wieder die Deter-

minantenfläche durch f und daran die Tangentialebene im Punkte f . Ist $g = (b_{hk})$ eine beliebige *andere positive Kantenform* dieses Raumes und cg dasjenige Multiplum von g , welches in jene Tangentialebene fällt, so darf beim Fortgang auf der geradlinigen Strecke von f nach cg hin das Minimum weder zunehmen noch konstant bleiben, d. h. während

$$c \sum \frac{\partial D(f)}{\partial a_{hk}} b_{hk} = nD(f)$$

ist, muß gleichzeitig $cb_{11} < M(f) = a_{11}$ sein, oder also *es muß*

$$(43) \quad \frac{1}{nD(f)} \sum \frac{\partial D(f)}{\partial a_{hk}} b_{hk} > \frac{b_{11}}{a_{11}}$$

sein.

5. *Erfüllt andererseits eine positive Kantenform f des reduzierten Raumes in bezug auf jede andere positive Kantenform g dieses Raumes die vorstehende Bedingung (43), so ist in der Tat f eine extreme Form in bezug auf den reduzierten Raum.*

Denn für jede nicht wesentlich positive Kantenform g des reduzierten Raumes gilt diese Ungleichung (43) ohne weiteres (vgl. die Ungleichung (40)); für f selbst an Stelle von g gilt die aus (43) durch Änderung des Zeichens $>$ in $=$ hervorgehende Beziehung; und nach der in (35) gefundenen Darstellung jeder beliebigen reduzierten Form als Aggregat $\sum cg$ aus Kantenformen würde dann diese Ungleichung (43) überhaupt für jede beliebige reduzierte Form g hervorgehen, die nicht ein bloßes Vielfaches von f ist.

In dieser Allgemeinheit aber besagt die Ungleichung (43) in der Tat, daß bei jeder infinitesimalen Variation von f , wobei die variierte Form im reduzierten Raume verbleibt und auch nicht bloß auf dem Strahle vom Nullpunkte aus durch f sich bewegt, die Größe $\frac{M(f)}{\sqrt[n]{D(f)}}$ stets abnimmt.

6. *Soll nun eine Klasse f eine extreme sein, so ist jedenfalls notwendig, daß jede einzelne in der Klasse vorhandene reduzierte Form eine extreme Form in bezug auf den reduzierten Raum ist.* Es gilt aber auch die Umkehrung hiervon und erlangen wir damit folgendes Charakteristikum der extremen Klassen:

Eine positive Formenklasse f ist nur dann und immer dann eine extreme Klasse, wenn jede einzelne reduzierte Form der Klasse eine extreme Form in bezug auf den reduzierten Raum ist.

In der Tat, es sei für die Klasse einer positiven Form $f = \sum a_{hk} x_h x_k$ die hier genannte Bedingung erfüllt. Wir bestimmen *jede existierende unimodulare ganzzahlige Substitution S , welche f in eine reduzierte Form überführt.* Wir bestimmen weiter jedesmal für die sich durch die Sub-

stitution S ergebende reduzierte Form $g = \sum b_{hk} y_h y_k$ eine positive GröÙe ε_S folgender Art: Das Gebiet aller Formen $g^* = \sum (b_{hk} + \varepsilon_{hk}) y_h y_k$, für welche alle Beträge $|\varepsilon_{hk}| < \varepsilon_S$ sind, soll, soweit es in den reduzierten Raum hineinfällt, dort nur solche Seitenwände dieses Raumes treffen, die auch g enthalten, und, außer auf der Kante vom Nullpunkte durch g , überall kleinere Werte der Funktion $M(f) : \sqrt[n]{D(f)}$ als im Punkte g darbieten.

Wir ermitteln endlich eine positive GröÙe δ derart, daß alle Formen $\sum \delta_{hk} x_h x_k$, wobei die Beträge $|\delta_{hk}| < \delta$ sind, durch die Substitutionen S nur in solche Formen $\sum \varepsilon_{hk} y_h y_k$ übergehen, in denen dann die Beträge $|\varepsilon_{hk}| < \varepsilon_S$ sind.

Alsdann kann der ganze Bereich der durch $f^* = \sum (a_{hk} + \delta_{hk}) x_h x_k$ mit den Bedingungen $|\delta_{hk}| < \delta$ dargestellten Formen in den einzelnen, mit dem reduzierten Raume B äquivalenten Kammern $B_{S^{-1}}$, denen f angehört, immer nur solche Wände treffen, die auch f enthalten, und kann daher nicht aus diesen Kammern $B_{S^{-1}}$ heraustreten. Es gibt daher zu jeder solchen Form f^* wenigstens eine von den Substitutionen S , welche sie gleichzeitig mit f in eine reduzierte Form überführt, und danach wird in diesem ganzen Bereiche, außer auf dem vom Nullpunkte aus durch f gehenden Strahl, die Funktion $M(f) : \sqrt[n]{D(f)}$ stets kleiner als in f sein.

7. Diejenigen positiven Kantenformen auf der Fläche $D(f) = 1$, welche den allergrößten Wert von a_{11} haben, repräsentieren notwendig extreme Klassen und bestimmen die präzise obere Grenze aller Werte von $M(f) : \sqrt[n]{D(f)}$.

§ 11. Die binären, ternären, quaternären Formen.

Auf Grund der Ergebnisse des § 5 und § 6 können wir für jeden Wert n die Seitenwände und die Kanten des reduzierten Raumes angeben und können wir nunmehr auch alle extremen Formenklassen ermitteln.

1. Man findet beispielsweise, daß in den Fällen $n = 2, 3, 4$ bereits diejenigen Ungleichungen

$$f(s_1, s_2, \dots, s_n) \geq a_{11} \quad (l = 1, 2, \dots, n),$$

worin $s_l = 1$ und die übrigen Zahlen $s_h = \pm 1$ oder zum Teil $= 0$ sind, alle übrigen der Ungleichungen (I) nach sich ziehen.

Nämlich aus diesen besonderen Ungleichungen läßt sich dann bereits allgemein

$$f(m_1, m_2, \dots, m_n) \geq a_{11}$$

für jedes beliebige System von ganzen Zahlen $m_1, m_2, \dots, m_n$, wobei $m_1, m_{l+1}, \dots, m_n$ nicht sämtlich Null sind, erschließen.

Da wir uns vorbehalten können, von den Substitutionen

$$x_1 = \pm y_1, \quad x_2 = \pm y_2, \quad \dots, \quad x_n = \pm y_n$$

Gebrauch zu machen, genügt es, die hiermit behauptete Tatsache für den Fall einzusehen, daß $m_1, m_2, \dots, m_n$ sämtlich ≥ 0 sind. Andererseits dürfen wir bei dem Nachweis dieser engeren Tatsache die entsprechenden Ungleichungen für die kleineren Werte des n bereits als sichergestellt annehmen, so daß wir überhaupt nur noch Werte $m_1, m_2, \dots, m_n$, die sämtlich > 0 sind, und dabei den Index $l = n$ in Betracht zu ziehen brauchen.

Unter den positiven Werten $m_1, m_2, \dots, m_n$ sei m_j der letzte von kleinstem Betrage. Wir setzen $u_h = m_j$, wenn $h \neq j$ ist, und $u_j = 0$, dabei wird immer

$$m_h - u_h \geq 0, \quad m_n - u_n > 0$$

sein, und die n Argumente $m_h - u_h$ erscheinen gegenüber den Argumenten m_h verringert bis auf eines, das unverändert geblieben ist. Nun haben wir:

$$\begin{aligned} f(m_1, m_2, \dots, m_n) - f(m_1 - u_1, m_2 - u_2, \dots, m_n - u_n) \\ = m_j^2 (f(1, 1, \dots, 1) - a_{jj}) + 2 \sum_{h \neq j} (m_h - m_j) m_j \sum_{k \neq j} a_{hk}, \end{aligned}$$

und die rechte Seite hier erweist sich durch $f(1, 1, \dots, 1) \geq a_{jj}$ und durch die Ungleichungen $a_{hh} + 2a_{hk} \geq 0$ in Anbetracht von $n \leq 4$ als ≥ 0 , so daß

$$f(m_1, m_2, \dots, m_n) \geq f(m_1 - u_1, m_2 - u_2, \dots, m_n - u_n)$$

folgt. Damit ist hier eine Rekursionsformel gewonnen, die durch einen Induktionsschluß sofort den verlangten Beweis liefert.

2. Stellen wir eine Form f durch das quadratische Schema ihrer Koeffizienten dar, so sind in den Fällen $n = 2, 3, 4$ die positiven Kantenformen mit $M(f) = 2$:

$$\begin{pmatrix} 2, 1 \\ 1, 2 \end{pmatrix}, \quad \begin{pmatrix} 2, 1, 1 \\ 1, 2, 1 \\ 1, 1, 2 \end{pmatrix}, \quad \begin{pmatrix} 2, 1, 1, 1 \\ 1, 2, 1, 1 \\ 1, 1, 2, 1 \\ 1, 1, 1, 2 \end{pmatrix} \quad \text{und} \quad \begin{pmatrix} 2, 0, 0, 1 \\ 0, 2, 0, 1 \\ 0, 0, 2, 1 \\ 1, 1, 1, 2 \end{pmatrix}$$

sowie die diesen äquivalenten reduzierten Formen; die bezüglichen Determinanten sind

$$3, 4, 5 \quad \text{und} \quad 4.$$

Alle diese Formen sind extreme Formen. Im vierdimensionalen Raume existieren danach zwei wesentlich verschiedene dichteste gitterförmige Lagerungen von lauter gleichen Kugeln.

Die nicht wesentlich positiven Kantenformen erhält man bei der Schreibweise hier aus den positiven Kantenformen für die geringeren

Variablenzahlen durch Vorsetzen so vieler aus lauter Nullen bestehender Horizontal- und Vertikalreihen, daß quadratische Systeme von n Reihen resultieren.

3. Auf die Fälle $n = 5$ und $n = 6$ bin ich in einem Aufsätze in Crelles Journal, Bd. 101 (diese Ges. Abhandlungen, Bd. I, S. 217) eingegangen. Für $n = 5$ haben Korkine und Zolotareff die extremen Formenklassen in den Mathematischen Annalen, Bd. 11 bestimmt.

§ 12. Volumen des reduzierten Raumes bis zur Determinantenfläche.

Unter dem *Volumen eines Bereichs in der Mannigfaltigkeit A* der quadratischen Formen verstehen wir den Wert des $\frac{n(n+1)}{2}$ -fachen Integrals

$$\iint \dots \int da_{11} da_{12} \dots da_{nn},$$

erstreckt über diesen Bereich.

1. Es sei D irgendein fester positiver Wert. Wir wollen den Satz beweisen:

Dasjenige Gebiet $B(D)$ des reduzierten Raumes, in welchem $D(f) \leq D$ gilt, welches also in bezug auf die Determinantenfläche $D(f) = D$ auf der Seite des Nullpunktes liegt, besitzt ein bestimmtes endliches Volumen.

Da die einzelnen Gebiete dieser Art, welche zu verschiedenen Werten D gehören, untereinander homothetisch vom Nullpunkte aus sind, so wird das fragliche Volumen einen Ausdruck $v_n D^{\frac{n+1}{2}}$ haben, wobei v_n eine nur von n abhängende Konstante sein wird.

2. Wir bezeichnen mit $B(D, \varepsilon)$ den Teil des reduzierten Raumes, worin

$$D(f) \leq D, \quad a_{11} \geq \varepsilon$$

gilt, unter ε eine positive Größe verstanden. Diese Ungleichungen in Verbindung mit den Ungleichungen

$$(44) \quad a_{11} \leq a_{22} \leq \dots \leq a_{nn}, \quad \pm 2a_{hk} \leq a_{hh} \quad (h < k),$$

$$(45) \quad \lambda_n a_{11} a_{22} \dots a_{nn} \leq D(f)$$

für eine reduzierte Form liefern *obere Grenzen für die Beträge aller Koordinaten in $B(D, \varepsilon)$* , und kommt daher diesem Bereiche $B(D, \varepsilon)$ ein bestimmtes endliches Volumen zu.

3. Ist nun $\varepsilon > \varepsilon^* > 0$, so umfaßt $B(D, \varepsilon^*)$ das Gebiet $B(D, \varepsilon)$ und in der Partie, mit welcher $B(D, \varepsilon^*)$ über $B(D, \varepsilon)$ hinausragt, gilt erstlich:

$$\varepsilon^* \leq a_{11} \leq \varepsilon, \quad |a_{1k}| \leq \frac{1}{2} a_{11} \quad (k = 2, 3, \dots, n), \quad a_{12} \geq 0,$$

$$\lambda_n a_{11} \frac{\partial D(f)}{\partial a_{11}} \leq D,$$

und ist darin zudem $\sum_2^n a_{hk} x_h x_k$ eine Form von $n-1$ Variablen mit der Determinante $\frac{\partial D(f)}{\partial a_{11}}$, welche alle Bedingungen einer reduzierten solchen Form erfüllt.

Nehmen wir nun das zu beweisende Resultat als bereits sichergestellt für den Fall von $n-1$ Variablen an, so vergrößert sich hiernach beim Übergang von $B(D, \varepsilon)$ zu $B(D, \varepsilon^*)$ das Volumen dieses Bereichs gewiß um weniger, als der Wert des Integrals

$$\frac{1}{2} \int a_{11}^{n-1} v_{n-1} \left(\frac{D}{\lambda_n a_{11}} \right)^{\frac{n}{2}} da_{11}$$

über den Bereich $0 < a_{11} \leq \varepsilon$ beträgt, d. i. um weniger als

$$\frac{v_{n-1}}{n \lambda_n^{\frac{n}{2}}} \varepsilon^{\frac{n}{2}} D^{\frac{n}{2}}.$$

Daraus folgt, daß das Volumen von $B(D, \varepsilon)$ mit nach Null abnehmendem ε einer bestimmten endlichen Grenze zustrebt, welche eben das Volumen von $B(D)$ definiert.

Nachdem so die Existenz der Konstante v_n erwiesen ist, können wir hinzusetzen:

Das Volumen von $B(D, \varepsilon)$ ist

$$(46) \quad < v_n D^{\frac{n+1}{2}} \quad \text{und} \quad > v_n D^{\frac{n+1}{2}} - \tilde{v}_n \varepsilon^{\frac{n}{2}} D^{\frac{n}{2}},$$

wo $\tilde{v}_n$ zur Abkürzung für $\frac{1}{n} v_{n-1} \lambda_n^{-\frac{n}{2}}$ steht.

4. Wir bemerken ferner:

Das durch

$$(47) \quad D \leq D(f) \leq D^*, \quad a_{11} \geq \varepsilon$$

definierte Gebiet des reduzierten Raumes, wobei $0 < D < D^*$ und $\varepsilon > 0$ sei, hat ein Volumen, das

$$(48) \quad < v_n \left(D^{*\frac{n+1}{2}} - D^{\frac{n+1}{2}} \right) \quad \text{und} \quad > \left(v_n - \tilde{v}_n \frac{\varepsilon^{\frac{n}{2}}}{D^{\frac{1}{2}}} \right) \left(D^{*\frac{n+1}{2}} - D^{\frac{n+1}{2}} \right)$$

ist.

Das Gebiet (47) nämlich ist ganz enthalten in dem Teile

$$D \leq D(f) \leq D^*$$

des reduzierten Raumes, woraus die obere Grenze in (48) hervorgeht.

Andererseits enthält jenes Gebiet ganz das Gebiet

$$D \leq D(f) \leq D^*, \quad \vartheta = \frac{\varepsilon}{\sqrt[n]{D}} \leq \frac{a_{11}}{\sqrt[n]{D(f)}}$$

des reduzierten Raumes. Das Volumen des letzteren Gebiets ist das $\left(D^{*\frac{n+1}{2}} - D^{\frac{n+1}{2}}\right) : D^{\frac{n+1}{2}}$ -fache des Volumens des Gebiets

$$(49) \quad D(f) \leq D, \quad \vartheta \leq \frac{a_{11}}{\sqrt[n]{D(f)}}$$

vom reduzierten Raume, da die durch die letzteren Bedingungen bei festem ϑ und verschiedenen Werten D definierten Gebiete homothetische Kegel vom Nullpunkte aus darstellen. Das durch (49) bestimmte Gebiet des reduzierten Raumes endlich enthält ganz das Gebiet $B(D, \varepsilon)$, woraus die untere in (48) genannte Grenze hervorgeht.

§ 13. Verwendung Dirichletscher Reihen.

Wir werden nunmehr zur Ermittlung des Volumens v_n wesentlich diejenigen Methoden heranziehen, auf welche Dirichlet die Bestimmung der Klassenanzahlen in der Theorie der binären Formen gegründet hat. Wir werden an Stelle des Volumenintegrals über den Bereich $B(D)$ das Integral einer gewissen Funktion über diesen Bereich betrachten, welche in einem überwiegenden Teile dieses Bereiches angenähert gleich 1 ist, und vermöge des Ausdrucks dieser Funktion wird eine Zurückführung der Bestimmung von v_n auf diejenige von v_{n-1} gelingen.

1. Es seien ε , G und σ positive Größen; wir werden schließlich unter Forderung eines gewissen Zusammenhanges unter ihnen G über jede Grenze wachsen, ε und σ nach Null abnehmen lassen. Wir setzen bereits $\sigma < \frac{1}{2}$, $G > \varepsilon$ und etwa $\varepsilon \leq 1$ voraus.

Wir haben es hier mit dem folgenden Ausdrucke zu tun:

$$(50) \quad \Phi(f) = \sigma(D(f))^{\frac{1}{2} + \frac{\sigma}{n}} \sum \frac{1}{(f(x_1, \dots, x_n))^{\frac{n}{2} + \sigma}}.$$

Darin bedeute $f(x_1, \dots, x_n) = \sum a_{hk} x_h x_k$ eine positive reduzierte quadratische Form von n Variablen, $D(f)$ ihre Determinante und die Summe soll über alle diejenigen Systeme von ganzen Zahlen $x_1, \dots, x_n$ erstreckt werden, welche die Ungleichungen

$$(51) \quad \varepsilon \leq f(x_1, \dots, x_n) < G$$

erfüllen und wobei $x_1, \dots, x_n$ keinen gemeinsamen Teiler > 1 haben.

2. Es bedeute andererseits $\Psi(f)$ den Wert des Ausdrucks (50), wenn darin die Summe ausnahmslos über alle existierenden ganzzahligen Systeme $x_1, \dots, x_n$ erstreckt wird, welche den Bedingungen (51) genügen. Wir lassen also bei der Definition von $\Psi(f)$ die Bedingung fallen, daß $x_1, \dots, x_n$ relativ prim sein sollen.

Wir erinnern noch an die in § 7, 1. gefundene Tatsache:

Sind $x_1, \dots, x_n$ ganze Zahlen $\neq 0, \dots, 0$ und ist darunter x_h die letzte von Null verschiedene Zahl, so fällt dafür

$$f(x_1, \dots, x_n) \geq \lambda_n a_{hh}$$

aus.

3. Die Untersuchung des Ausdrucks $\Psi(f)$ gründen wir auf folgende Tatsachen.

Deuten wir $x_1, \dots, x_n$ als Koordinaten eines Punktes in einem n -dimensionalen Raume $\mathfrak{R}_n$, so stellt

$$(52) \quad f(x_1, \dots, x_n) < T,$$

wenn T eine positive Konstante ist, das Innere eines n -dimensionalen Ellipsoids in diesem Raume vor. Das Volumen in $x_1, \dots, x_n$ hat für dieses Ellipsoid den Wert

$$(53) \quad \gamma_n \frac{(\sqrt{T})^n}{\sqrt{D(f)}},$$

wobei

$$(54) \quad \gamma_n = \frac{\pi^{\frac{n}{2}}}{\Gamma(1 + \frac{n}{2})} = \frac{\pi^{\left[\frac{n}{2}\right]}}{\frac{n}{2} \left(\frac{n}{2} - 1\right) \dots \left(\frac{n}{2} - \left[\frac{n-1}{2}\right]\right)}$$

das Volumen einer n -dimensionalen Kugel vom Radius 1 vorstellt (vgl. § 9).

Wir bemerken ferner, daß der Würfelbereich

$$(55) \quad -\frac{1}{2} \leq x_1 \leq \frac{1}{2}, \dots, -\frac{1}{2} \leq x_n \leq \frac{1}{2}$$

den Ungleichungen (44) zufolge völlig in das Ellipsoid

$$(56) \quad f(x_1, \dots, x_n) \leq \frac{n(n+1)}{8} a_{nn}$$

zu liegen kommt.

Auf Grund dieser Umstände erhalten wir eine approximative Bestimmung für die Anzahl derjenigen ganzzahligen Systeme (Gitterpunkte) $x_1, \dots, x_n$, welche die Bedingung (52) erfüllen.

Wir konstruieren nämlich um jeden einzelnen der betreffenden Gitterpunkte als Mittelpunkt einen Würfel mit Seitenflächen parallel den Koordinatenebenen und von der Kantenlänge 1, d. i. jedesmal den Würfel, der durch Parallelverschiebung des Würfels (55) vom Nullpunkte nach dem betreffenden Gitterpunkte entsteht.

Da wir jeden dieser Würfel durch ein dem Ellipsoide (56) homologes Ellipsoid umschließen können, fällt der gesamte Bereich dieser Würfel völlig in das Gebiet des Ellipsoids

$$f(x_1, \dots, x_n) < \left(\sqrt{T} + \sqrt{\frac{n(n+1)}{8} a_{nn}} \right)^2$$

hinein. Alle jene Würfel dringen nicht gegenseitig ineinander ein und sind je vom Volumen 1. Danach ist die Anzahl der fraglichen Gitterpunkte sicher

$$< \frac{\gamma_n}{\sqrt{D(f)}} \left(\sqrt{T} + \sqrt{\frac{n(n+1)}{8} a_{nn}} \right)^n.$$

Wir schreiben zur Abkürzung

$$\gamma_n \left\{ \left(1 + \sqrt{\frac{n(n+1)}{8 \lambda_n}} \right)^n - 1 \right\} = \kappa_n,$$

und wir können behaupten:

Die Anzahl der ganzzahligen Auflösungen von

$$f(x_1, \dots, x_n) < T$$

ist, wenn $T \geq \lambda_n a_{nn}$ ist, sicher

$$(57) \quad < \frac{1}{\sqrt{D(f)}} \left(\gamma_n T^{\frac{n}{2}} + \kappa_n (\lambda_n a_{nn})^{\frac{1}{2}} T^{\frac{n-1}{2}} \right).$$

Andererseits überzeugen wir uns davon, daß die vorhin konstruierten Würfel, falls $T > \frac{n(n+1)}{8} a_{nn}$ ist, gewiß das Gebiet des kleineren Ellipsoids

$$f(x_1, \dots, x_n) < \left(\sqrt{T} - \sqrt{\frac{n(n+1)}{8} a_{nn}} \right)^2$$

völlig überlagern, und hieraus leiten wir die folgende andere Tatsache her:

Die Anzahl der ganzzahligen Auflösungen von

$$f(x_1, \dots, x_n) < T$$

ist, wenn $T \geq \lambda_n a_{nn}$ ist, sicher

$$(58) \quad > \frac{1}{\sqrt{D(f)}} \left(\gamma_n T^{\frac{n}{2}} - \kappa_n (\lambda_n a_{nn})^{\frac{1}{2}} T^{\frac{n-1}{2}} \right).$$

Aus den beiden in (57) und (58) enthaltenen Grenzen entnehmen wir weiter:

Ist $T \geq \lambda_n a_{nn}$ und t ein Wert > 1 , so liegt die Anzahl der ganzzahligen Systeme $x_1, \dots, x_n$, welche die Ungleichungen

$$T \leq f(x_1, \dots, x_n) < Tt$$

befriedigen, jedenfalls innerhalb der zwei Grenzen

$$(59) \quad \frac{1}{\sqrt{D(f)}} \left(\gamma_n (t^{\frac{n}{2}} - 1) T^{\frac{n}{2}} \pm \kappa_n (t^{\frac{n-1}{2}} + 1) (\lambda_n a_{nn})^{\frac{1}{2}} T^{\frac{n-1}{2}} \right).$$

4. Wir fassen zunächst diejenigen Glieder aus $\Psi(f)$ zusammen, in denen $f \geq \lambda_n a_{nn}$ und dabei jedenfalls auch $f \geq \varepsilon$ ist. Es sei T_n die größere der zwei unteren Schranken ε und $\lambda_n a_{nn}$ (bzw. ihr gemeinsamer Wert). Auf die letzte Angabe gestützt, können wir das ganze Aggregat der betreffenden Glieder aus $\Psi(f)$ sofort in zwei Grenzen einschließen.

Es sei $G > \lambda_n a_{nn}$. Wir setzen $G = T_n t^l$, so daß der Exponent l eine positive ganze Zahl und der Wert $t > 1$ wird, und wir teilen das Intervall $T_n \leq f < G$ durch Einschalten von $T_n t, \dots, T_n t^{l-1}$ in die l Intervalle

$$T_n \leq f < T_n t, \dots, T_n t^{l-1} \leq f < G.$$

Wir bestimmen für die einzelnen Intervalle nach (59) eine obere (bzw. untere) Grenze der Anzahl der ganzzahligen Systeme $x_1, \dots, x_n$, für die f in das Intervall fällt, und ersetzen in den Nennern der betreffenden Glieder aus $\Psi(f)$ die Größe $f(x_1, \dots, x_n)$ durch die untere (bzw. obere) Grenze des Intervalls. Dadurch erhalten wir eine obere (bzw. untere) Grenze für jenes Aggregat aus $\Psi(f)$.

Wir finden zunächst, daß jener Anteil aus $\Psi(f)$

$$(60) < \left\{ \gamma_n \frac{\sigma(t^{\frac{n}{2}} - 1)}{\left(1 - \frac{1}{t^\sigma}\right)} \left(\frac{1}{T^\sigma} - \frac{1}{G^\sigma}\right) + \kappa_n \left(t^{\frac{n-1}{2}} + 1\right) \frac{\sigma}{1 - \frac{1}{t^{\frac{1}{2} + \sigma}}} T_n^{\frac{1}{2}} \left(\frac{1}{T_n^{\frac{1}{2} + \sigma}} - \frac{1}{G^{\frac{1}{2} + \sigma}}\right) \right\} (D(f))^{\frac{\sigma}{n}}$$

ist.

5. Es sei D eine feste positive Größe. Wir denken uns *augenblicklich* f speziell im Bereiche $B(D, \varepsilon)$ gelegen. Dann ist also $a_{11} \geq \varepsilon$ und aus (44) und (45) folgt

$$\lambda_n \varepsilon^n \leq D(f) \leq D, \quad \lambda_n \varepsilon \leq \lambda_n a_{nn} \leq \frac{D}{\varepsilon^{n-1}}.$$

Wir verfügen bei dem beabsichtigten Grenzprozeß

$$(61) \quad \lim \varepsilon = 0, \quad \lim \sigma = 0, \quad \lim G = \infty$$

über σ, ε, G derart, daß dabei

$$(62) \quad \lim \varepsilon^\sigma = 1, \quad \lim G^{-\sigma} = 0$$

wird. Alsdann wird auch $\lim T_n^\sigma = 1$ sein.

Ferner können wir l als ganze Zahl abhängig von ε, σ und G noch derart einrichten, daß hierbei

$$\frac{1}{\log t} = \frac{\sigma l}{\sigma(\log G - \log T_n)}$$

nach unendlich, aber $\frac{\sigma}{\log t}$ nach Null konvergiert. Dann konvergiert also $\log t$ nach Null, mithin t nach 1. Infolge dieser Umstände konvergiert der Term mit κ_n in (60) nach Null. In dem ersten Term mit dem Koeffizienten γ_n ist der Faktor

$$\frac{\sigma(t^{\frac{n}{2}} - 1)}{\left(1 - \frac{1}{t^\sigma}\right)} = \frac{n}{2} t^\sigma t^{\frac{n}{2} - \sigma},$$

wo t^* einen Mittelwert zwischen 1 und t bedeutet, und konvergiert dieser erste Term in (60) nach $\frac{n}{2} \gamma_n$.

Die analog herzustellende untere Grenze des fraglichen Anteils aus $\Psi(f)$ wird von dem Ausdrucke (60) her dadurch gewonnen, daß der zweite Term subtraktiv statt additiv genommen und beiden Termen noch der Faktor $t^{-\frac{n}{2}-\sigma}$ hinzugesetzt wird. Es leuchtet ein, daß unter den angegebenen Voraussetzungen diese untere Grenze ebenfalls nach dem Werte $\frac{n}{2} \gamma_n$ konvergiert.

6. Wir haben weiter diejenigen Terme des Ausdrucks $\Psi(f)$ abzuschätzen, in welchen f zwar $\geq \varepsilon$, aber $< \lambda_n a_{nn}$ ausfällt. Da die Konstante $\lambda_n < 1$ und a_{11} das Minimum von f ist, wird in allen Gliedern von $\Psi(f)$ stets $f > \lambda_n a_{11}$ sein.

Wir nehmen alle diejenigen Glieder aus $\Psi(f)$ zusammen, soweit solche überhaupt vorhanden sind, in welchen

$$\left. \begin{array}{l} \lambda_n a_{hh} \leq \\ \varepsilon \leq \end{array} \right\} f(x_1, \dots, x_n) < \lambda_n a_{h+1, h+1}$$

gilt, wobei h eine der Zahlen $1, 2, \dots, n-1$ sein kann. Es sei T_h die größere der zwei Größen $\lambda_n a_{hh}$, ε bzw. ihr gemeinsamer Wert. Wegen $f < \lambda_n a_{h+1, h+1}$ müssen hierbei notwendig $x_{h+1}, \dots, x_n$ sämtlich Null sein; das ganze Aggregat der in Rede stehenden Glieder ist daher sicherlich kleiner als der Wert der unendlichen Reihe

$$\sigma(D(f))^{\frac{1}{2} + \frac{\sigma}{n}} \sum \frac{1}{(f(x_1, \dots, x_h, 0, \dots, 0))^{\frac{n}{2} + \sigma}},$$

erstreckt über alle vorhandenen ganzzahligen Systeme $x_1, \dots, x_h$, bei welchen

$$T_h \leq f(x_1, \dots, x_h, 0, \dots, 0)$$

ist.

Wir übertragen das in der Formel (59) entwickelte Resultat auf den Fall von Formen mit h Variablen, und wir finden die Anzahl derjenigen ganzzahligen Systeme $x_1, \dots, x_h$, für die eine Einschränkung

$$T_h t^j \leq f(x_1, \dots, x_h, 0, \dots, 0) < T_h t^{j+1} \quad (j \geq 0)$$

statthat, unter t eine fest angenommene positive Konstante (etwa 2) verstanden, kleiner als eine Größe

$$\bar{\gamma}_h \frac{(T_h t^j)^{\frac{h}{2}}}{\sqrt{a_{11} a_{22} \dots a_{hh}}},$$

worin $\bar{\gamma}_h$ eine gewisse nur von h abhängende positive Konstante vorstellt. Das Aggregat der in Rede stehenden Terme ist dann sicher

$$< \sigma(D(f))^{\frac{1}{2} + \frac{\sigma}{n}} \frac{\bar{\gamma}_h}{\left(1 - \frac{1}{t^{\frac{n-h}{2} + \sigma}}\right) T_h^{\frac{n-h}{2} + \sigma} \sqrt{a_{11} a_{22} \dots a_{hh}}}.$$

Hierin machen wir im Nenner von

$$\lambda_n^{-\frac{n-h-1}{2}} T_h^{\frac{n-h}{2} + \sigma} \sqrt{a_{11} a_{22} \dots a_{hh}} \geq \varepsilon^{\frac{1}{2} + \sigma} a_{hh}^{\frac{n-h-1}{2}} \sqrt{a_{11} a_{22} \dots a_{hh}} \geq \varepsilon^{\frac{1}{2} + \sigma} a_{11}^{\frac{n-1}{2}}$$

Gebrauch, und mit Rücksicht hierauf kommen wir sogleich zu dem Ergebnisse:

Das Aggregat aller Terme in $\Psi(f)$, welche den Bedingungen

$$\varepsilon \leq f(x_1, \dots, x_n) < \lambda_n a_{nn}$$

entsprechen, ist

$$(63) \quad < \mu_n \frac{\sigma(D(f))^{\frac{1}{2} + \frac{\sigma}{n}}}{\varepsilon^{\frac{1}{2} + \sigma} a_{11}^{\frac{n-1}{2}}},$$

worin μ_n eine gewisse nur von n abhängende Konstante bedeutet.

7. Wir setzen jetzt wieder die Form f speziell im Gebiete $B(D, \varepsilon)$ gelegen voraus, so daß $D(f) \leq D$, $a_{11} \geq \varepsilon$ ist. Wir verfügen über σ im Zusammenhang mit ε noch so, daß für $\lim \varepsilon = 0$ auch

$$\lim \left(\frac{\sigma}{\varepsilon^{\frac{n}{2}}} \right) = 0$$

wird. Dabei wird von selber auch

$$\log(\varepsilon^\sigma) = \left(\sigma \varepsilon^{-\frac{n}{2}} \right) \left(\frac{n}{\varepsilon^{\frac{n}{2}}} \log \varepsilon \right)$$

nach Null, also ε^σ nach 1 konvergieren. Der vorstehende Ausdruck (63) wird nunmehr für die Form f nach Null konvergieren, und wir gelangen zu dem Resultate:

Liegt f im Gebiete $B(D, \varepsilon)$, so läßt sich der Wert von $\Psi(f)$ in zwei Grenzen einschließen, die unter den Voraussetzungen

$$(64) \quad \lim \varepsilon = 0, \lim \left(\frac{\sigma}{\varepsilon^{\frac{n}{2}}} \right) = 0, \lim G^{-\sigma} = 0$$

beide zugleich nach dem Werte $\frac{n}{2} \gamma_n$ konvergieren.

8. Es sei immer noch f speziell in $B(D, \varepsilon)$ gelegen. Um von dem Ausdrucke $\Psi(f)$ zu dem für uns eigentlich notwendigen Ausdrucke $\Phi(f)$ zu kommen, wollen wir mit Ψ_d den Wert des Ausdrucks (50) bezeichnen, wenn darin die Summe über alle solchen, den Bedingungen $\varepsilon \leq f < G$ entsprechenden ganzzahligen Systeme $x_1, \dots, x_n$ erstreckt wird, bei denen $x_1, \dots, x_n$ sämtlich durch eine bestimmte positive Zahl d teilbar sind. Bei Werten d , für welche $d^2 \varepsilon \geq G$ ist, wird es Systeme $x_1, \dots, x_n$ der eben bezeichneten Art überhaupt nicht geben und ist daher $\Psi_d = 0$ zu setzen. Die vorhin behandelte Summe $\Psi(f)$ kommt auf die Summe Ψ_1 hinaus.

Wir gelangen nun von Ψ_1 zu der von uns gesuchten Summe $\Phi(f)$, indem wir, der Reihe der Primzahlen 2, 3, 5, ... folgend, aus dem Aggregate Ψ_1 sukzessive alle diejenigen Terme aussondern, in denen $x_1, \dots, x_n$ sämtlich durch 2, oder doch sämtlich durch 3, oder doch sämtlich durch 5, usf. aufgehen.*)

Danach ergibt sich

$$(65) \quad \Phi(f) = \Psi_1 - \Psi_2 - (\Psi_3 - \Psi_6) - (\Psi_5 - \Psi_{10} - \Psi_{15} + \Psi_{30}) - \dots$$

Die Bildung dieser Reihe ist am besten dahin zu beschreiben, daß das über alle Primzahlen $p = 2, 3, 5, \dots$ zu erstreckende unendliche Produkt

$$(66) \quad \prod \left(1 - \frac{1}{p^s}\right) = 1 - \frac{1}{2^s} - \left(\frac{1}{3^s} - \frac{1}{6^s}\right) - \left(\frac{1}{5^s} - \frac{1}{10^s} - \frac{1}{15^s} + \frac{1}{30^s}\right) - \dots$$

entwickelt und jedem in dieser Entwicklung auftretenden Gliede $\pm \frac{1}{d^s}$ ein Glied $\pm \Psi_d$ mit dem nämlichen Vorzeichen $\pm$ in der Reihe (65) zugeordnet wird.

Es sei hier $s = n + 2\sigma$, so ist die Reihe (66) absolut konvergent und ihr Wert das Reziproke der über alle positiven ganzen Zahlen erstreckten Summe

$$(67) \quad S_s = 1 + \frac{1}{2^s} + \frac{1}{3^s} + \frac{1}{4^s} + \dots$$

Setzen wir in den Gliedern von Ψ_d die Variablen $x_h = dy_h$ an, so wird darin $f(x_1, \dots, x_n) = d^{2s} f(y_1, \dots, y_n)$, und da f in $B(D, \varepsilon)$ liegen soll, das Minimum von f also nach Voraussetzung $\geq \varepsilon$ ist, so haben wir in Ψ_d die Summation über alle ganzzahligen Systeme $y_1, \dots, y_n$ zu erstrecken, für welche

$$\varepsilon \leq f(y_1, \dots, y_n) < \frac{G}{d^2}$$

wird.

Wir denken uns nun d^* als ganze Zahl, abhängig von σ , derart gewählt, daß die Summe der Beträge aller derjenigen Glieder in (66), welche Werten $d > d^*$ entsprechen, beliebig nahe an Null liegt und daß andererseits unter den Voraussetzungen (64) auch $\lim(d^{2\sigma}) = 1$ wird. Alsdann haben wir nach dem Ergebnisse in 7., wenn $d \leq d^*$ ist, unter den Voraussetzungen (64) jedenfalls

$$\lim \left(\Psi_d : \frac{1}{d^{n+2\sigma}} \right) = \frac{n}{2} \gamma_n.$$

Wenn aber $d > d^*$ ist, so gilt zum mindesten die Relation

$$0 \leq \Psi_d \leq \frac{1}{d^{n+2\sigma}} \Psi_1.$$

*) Vgl. hierzu Lipschitz, Über die asymptotischen Gesetze von gewissen Gattungen zahlentheoretischer Funktionen, Monatsbericht der Berliner Akademie, 1865, S. 174.

Mit Rücksicht auf diese Umstände folgt offenbar unter jenen Voraussetzungen

$$(68) \quad \lim \Phi(f) = \frac{n}{2} \gamma_n \lim \prod \left(1 - \frac{1}{p^{n+2\sigma}}\right) = \frac{n}{2} \frac{\gamma_n}{S_n}.$$

Damit sind wir zu folgendem Ergebnisse gelangt:

Für alle Formen f im ganzen Bereiche $B(D, \varepsilon)$ liegt der Wert der Funktion $\Phi(f)$ zwischen zwei Grenzen, die unter den Voraussetzungen (64) beide zugleich nach dem Werte $\frac{n}{2} \frac{\gamma_n}{S_n}$ konvergieren.

9. Es habe jetzt das Minimum a_{11} von f einen beliebigen Wert, so können wir für das Aggregat derjenigen Glieder in $\Psi(f)$, bei welchen außer $\varepsilon \leq f$ noch $\lambda_n a_{nn} \leq f$ gilt, durch (60) eine Konstante v_n als obere Grenze bestimmen, und für das Aggregat der übrigen Glieder in $\Psi(f)$ besteht die obere Grenze (63). Diese oberen Grenzen gelten um so mehr für den Wert der Summe $\Phi(f)$, und wir finden:

Es ist stets

$$(69) \quad 0 \leq \Phi(f) < \mu_n \frac{\sigma(D(f))^{\frac{1}{2} + \frac{\sigma}{n}}}{\varepsilon^{\frac{1}{2} + \sigma} a_{11}^{\frac{n-1}{2}}} + v_n,$$

worin μ_n und v_n zwei positive, nur von n abhängende Konstanten vorstellen.

10. Es sei jetzt δ eine positive GröÙe $\leq \varepsilon$ und wir betrachten das $\frac{n(n+1)}{2}$ -fache Integral

$$(70) \quad J(\delta) = \int \int \dots \int \Phi(f) da_{11} da_{12} \dots da_{nn}$$

der in (50) definierten Funktion $\Phi(f)$, erstreckt über den durch

$$D(f) \leq D, \quad a_{11} \geq \delta$$

bestimmten Teil $B(D, \delta)$ des reduzierten Raumes.

Nehmen wir zunächst $\delta = \varepsilon$, so ergibt das Resultat aus 8. in Verbindung mit den Sätzen aus § 12, daß unter den Voraussetzungen (64) jedenfalls

$$\lim J(\varepsilon) = \frac{n}{2} \frac{\gamma_n}{S_n} v_n D^{\frac{n+1}{2}}$$

ist.

Nun sei $\delta < \varepsilon$, so benutzen wir dazu noch in dem Gebiete, mit welchem der Bereich $B(D, \delta)$ über $B(D, \varepsilon)$ hinausragt, für die Funktion $\Phi(f)$ die obere Grenze (69). Das Volumen jenes Zusatzgebietes ist nach § 12 sicher $< \bar{v}_n \varepsilon^{\frac{n}{2}} D^{\frac{n}{2}}$. Im ersten Term von (69) ersetzen wir $(D(f))^{\frac{\sigma}{n}}$ durch die obere Grenze $D^{\frac{\sigma}{n}}$; andererseits ermitteln wir eine obere Grenze für das $\frac{n(n+1)}{2}$ -fache Integral

$$\iint \dots \int \frac{(D(f))^{\frac{1}{2}}}{a_{11}^{\frac{n-1}{2}}} da_{11} da_{12} \dots da_{nn},$$

erstreckt über den ganzen Bereich $B(D, \delta)$, indem wir ähnlich wie in § 12 vorgehen. Indem wir die Integration nach $a_{22}, a_{23}, \dots, a_{nn}$ und zwar zunächst über die Flächen $\frac{\partial D(f)}{\partial a_{11}} = \text{konst.}$ ausführen, ferner nach $a_{12}, \dots, a_{1n}$ integrieren, finden wir das betreffende Integral (vgl. § 12, 3.)

$$< \int \frac{a_{11}^{\frac{1}{2}}}{a_{11}^{\frac{n-1}{2}}} \frac{1}{2} a_{11}^{n-1} da_{11} \int C^{\frac{1}{2}} v_{n-1} d\left(C^{\frac{n}{2}}\right),$$

über den Bereich

$$0 < \lambda_n a_{11}^n \leq D, \quad 0 < \lambda_n a_{11} C \leq D$$

erstreckt. Das Doppelintegral hier wird

$$\frac{n}{n+1} v_{n-1} \left(\frac{D}{\lambda_n}\right)^{\frac{n+1}{2}} \int \frac{1}{2} a_{11}^{-\frac{1}{2}} da_{11} = \frac{n}{n+1} v_{n-1} \left(\frac{D}{\lambda_n}\right)^{\frac{n+1}{2} + \frac{1}{2n}}$$

Berücksichtigen wir nun, daß im ersten Term von (69) noch der Faktor $\frac{\sigma}{\frac{1}{2} + \sigma}$ auftritt, der unter den Voraussetzungen (64) nach Null kon-

vergiert, so können wir endlich den Satz aussprechen:

Das Integral $J(\delta)$, über den Bereich $B(D, \delta)$ erstreckt, wobei δ beliebig $\leq \varepsilon$ ist, liegt zwischen zwei Grenzen, die unter den Voraussetzungen (64) beide nach dem Werte

$$(71) \quad \frac{n}{2} \frac{\gamma_n}{S_n} v_n D^{\frac{n+1}{2}}$$

konvergieren.

§ 14. Auswertung des Volumens.

Auf Grund dieses Satzes können wir jetzt die Berechnung von v_n auf die Berechnung von v_{n-1} zurückführen.

1. Sind $p_1, p_2, \dots, p_n$ n ganze Zahlen ohne gemeinsamen Teiler > 1 , so kann man stets dazu ganzzahlige unimodulare Substitutionen P mit diesen Zahlen als erster Vertikalreihe finden. Ist P_0 eine erste solche Substitution, P eine beliebige Substitution dieser Art, so hat die Substitution $P_0^{-1}P$ als erste Vertikalreihe die Zahlen $1, 0, \dots, 0$ wie die identische Substitution, und können wir jedesmal in bestimmter Weise $P = P_0 QR$ setzen, so daß Q die erste Variable völlig ungeändert läßt, also in Q

auch noch die erste Horizontalreihe $1, 0, \dots, 0$ ist und R eine Substitution von der Gestalt

$$y_1 = z_1 + r_2 z_2 + \dots + r_n z_n, \quad y_2 = z_2, \dots, \quad y_n = z_n$$

wird.

2. Nun mögen ε , G , σ und $\delta \leq \varepsilon$ wie in § 13 vorausgesetzt sein, und f sei eine Form im Bereiche $B(D, \delta)$. Sodann seien $p_1, p_2, \dots, p_n$ solche ganzen Zahlen ohne gemeinsamen Teiler > 1 , daß

$$(72) \quad \varepsilon \leq f(p_1, p_2, \dots, p_n) = b_{11} < G$$

wird. Wir führen P_0, P, Q, R wie soeben ein. Die durch $P = P_0 Q R$ aus f hervorgehende Form hat b_{11} als ersten Koeffizienten. Wir schreiben diese Form

$$(73) \quad \varphi = \sum b_{hk} y_h y_k = b_{11} (y_1 + \frac{b_{12}}{b_{11}} y_2 + \dots + \frac{b_{1n}}{b_{11}} y_n)^2 + \psi,$$

$$(74) \quad \psi = c_{22} y_2^2 + 2c_{23} y_2 y_3 + \dots + c_{nn} y_n^2,$$

wobei

$$(75) \quad c_{hk} = b_{hk} - \frac{b_{1h} b_{1k}}{b_{11}} \quad (h \leq k; h, k = 2, 3, \dots, n)$$

gesetzt ist. Die Determinante von ψ wird dabei

$$= \frac{D(f)}{b_{11}}.$$

Wir können zunächst P_0 irgendwie unimodular mit $p_1, p_2, \dots, p_n$ als erster Vertikalreihe gewählt annehmen, sodann Q derart bestimmen, daß die Form $\psi(y_2, y_3, \dots, y_n)$ von $n - 1$ Variablen als solche reduziert wird, also gewisse nach § 5 und § 6 in endlicher Anzahl anzuweisende lineare Ungleichungen

$$(76) \quad \sum' m_{hk} c_{hk} \geq 0 \quad (h, k = 2, 3, \dots, n)$$

erfüllt, wobei im allgemeinen für Q die Wahl zwischen zwei entgegengesetzten Substitutionen $Q^*, -Q^*$ sein wird. Weiter können wir R so bestimmen, daß alle Ungleichungen

$$(77) \quad \pm \frac{b_{12}}{b_{11}} \leq \frac{1}{2}, \dots, \pm \frac{b_{1n}}{b_{11}} \leq \frac{1}{2}$$

eintreten, und endlich können wir uns noch zwischen Q^* und $-Q^*$ derart entscheiden, daß

$$(78) \quad b_{12} \geq 0$$

wird. Durch die vorstehenden Forderungen sind dann Q, R und damit auch $P = P_0 Q R$ in dem Falle eindeutig bestimmt, wo in keiner der dabei zu betrachtenden Ungleichungen (76), (77), (78) sich das Gleichheitszeichen einstellt.

3. Jetzt sei ϑ ein positiver Wert $\leq \delta$ und wir fassen für φ den Bereich aller derjenigen Formen ins Auge, bei welchen

$$(79) \quad \varepsilon \leq b_{11} < G, \quad D(\varphi) \leq D, \quad \vartheta \leq c_{22}$$

ist und zudem alle Ungleichungen (76), (77), (78) bestehen.

Das Minimum einer nach (73), (74) dargestellten Form φ ist gewiß nicht größer als das Minimum der binären Form

$$b_{11}(y_1 + \frac{b_{12}}{b_{11}}y_2)^2 + c_{22}y_2^2$$

und letzteres Minimum ist (vgl. § 5, 1.) $\leq \sqrt{\frac{4}{3}b_{11}c_{22}}$. Soll φ einer Form f in $B(D, \delta)$ äquivalent sein und (72) gelten, so ist daher notwendig

$$\delta \leq \sqrt{\frac{4}{3}} G c_{22},$$

und φ kommt sicher in den obigen Bereich zu liegen, wofern wir $\vartheta = \frac{3\delta^2}{4G}$ voraussetzen.

Andererseits ist eine beliebige durch φ mittels ganzer, nicht sämtlich verschwindender Zahlen $y_1, y_2, \dots, y_n$ darstellbare Größe entweder $\geq c_{22} \geq \vartheta$, weil c_{22} das Minimum von ψ ist, oder, wofern dabei $y_2, \dots, y_n$ sämtlich Null sind, doch $\geq b_{11} \geq \varepsilon \geq \delta \geq \vartheta$. Danach fallen die zu den Formen φ des obigen Bereichs äquivalenten reduzierten Formen f sicher stets in $B(D, \vartheta)$ hinein.

Für die Koeffizienten $b_{11}, b_{22}, \dots, b_{nn}$ in φ sind gewisse obere, bloß von $\varepsilon, G, \vartheta, D$ abhängende Grenzen angebar, und kommen dadurch bei gegebenen Werten dieser Größen von vornherein nur eine endliche Anzahl unimodularer ganzzahliger Substitutionen P in Frage, durch welche eine reduzierte Form f in $B(D, \vartheta)$ in eine den Bedingungen (76), (77), (78), (79) entsprechende Form φ übergehen könnte.

Danach verteilt sich der obige Bereich der Formen φ ganz auf eine endliche Anzahl der mit dem reduzierten Raume B äquivalenten Kammern B_p . Angenommen ferner, φ fällt weder auf die begrenzenden Seitenwände dieser Kammern noch auf eine der durch die Gleichheitszeichen in (76), (77), (78) angewiesenen Flächen; dann ist einerseits die zu φ äquivalente reduzierte Form f sowie das Paar entgegengesetzter ganzzahliger unimodularer Substitutionen $P, -P$, durch welche f in φ übergeht, eindeutig bestimmt; und andererseits gibt es auch keine von P und $-P$ verschiedene ganzzahlige unimodulare Substitution P^* mit der nämlichen ersten Vertikalreihe wie P oder $-P$, durch welche f in eine andere, ebenfalls allen jenen Ungleichungen (76), (77), (78), (79) genügende Form φ^* überginge.

4. Wir bilden nun das $\frac{n(n+1)}{2}$ -fache Integral

$$(80) \quad K(\vartheta) = \iint \cdots \int 2\sigma \frac{(D(\varphi))^{\frac{1}{2} + \frac{\sigma}{n}}}{b_{11}^{\frac{n}{2} + \sigma}} db_{11} db_{12} \cdots db_{nn},$$

erstreckt über den ganzen, durch die Ungleichungen (76), (77), (78), (79) definierten Bereich. Aus den soeben entwickelten Tatsachen geht dann folgendes Verhältnis dieses Integrals zu dem in (70) eingeführten Integrale $J(\delta)$ hervor:

Wir haben, wenn $\delta \leq \varepsilon$ ist:

$$(81) \quad J(\delta) < K\left(\frac{3\delta^2}{4G}\right) < J\left(\frac{3\delta^2}{4G}\right).$$

Um das Integral $K(\vartheta)$ auszuwerten, führen wir zunächst $c_{22}, c_{23}, \dots, c_{nn}$ anstatt $b_{22}, b_{23}, \dots, b_{nn}$ gemäß (75) ein, wobei die zugehörige Funktionaldeterminante den Wert 1 hat, sodann bewerkstelligen wir die Integration nach $c_{22}, c_{23}, \dots, c_{nn}$ und zwar zunächst über Flächen konstanter Determinante von ψ und hernach für alle in Betracht kommenden Werte C dieser Determinante, endlich leisten wir auch die Integration nach $b_{12}, \dots, b_{1n}$.

Übertragen wir das in § 12 (47) und (48) dargestellte Ergebnis auf den Fall von $n-1$ Variablen, so erhalten wir auf diese Weise als eine obere Grenze für $K(\vartheta)$ den Wert

$$\int 2\sigma b_{11}^{\frac{1}{2} + \frac{\sigma}{n}} \frac{1}{2} \frac{b_{11}^{n-1} db_{11}}{b_{11}^{\frac{n}{2} + \sigma}} \int C^{\frac{1}{2} + \frac{\sigma}{n}} v_{n-1} d\left(C^{\frac{n}{2}}\right),$$

über den Bereich

$$\varepsilon \leq b_{11} < G, \quad 0 < b_{11} C \leq D$$

erstreckt, d. i.

$$(82) \quad \frac{n}{n+1 + \frac{2\sigma}{n}} v_{n-1} D^{\frac{n+1}{2} + \frac{\sigma}{n}} \left(\frac{1}{\varepsilon^\sigma} - \frac{1}{G^\sigma} \right).$$

Sodann haben wir eine untere Grenze für das Integral $K(\vartheta)$, welche aus dieser oberen Grenze durch Verminderung um

$$\int \sigma b_{11}^{\frac{n}{2} - \frac{1}{2} - \sigma + \frac{\sigma}{n}} db_{11} \int C^{\frac{\sigma}{n}} \vartheta^{\frac{n-1}{2}} \tilde{v}_{n-1} d\left(C^{\frac{n}{2}}\right)$$

entsteht; letzterer Ausdruck ist

$$(83) \quad = \frac{n}{n + \frac{2\sigma}{n}} \tilde{v}_{n-1} D^{\frac{n}{2} + \frac{\sigma}{n}} \frac{\sigma}{\frac{1}{2} - \sigma} \vartheta^{\frac{n-1}{2}} \left(G^{\frac{1}{2} - \sigma} - \varepsilon^{\frac{1}{2} - \sigma} \right)$$

und konvergiert für $\vartheta \leq \frac{3\varepsilon^2}{4G}$ und unter den Voraussetzungen (64) nach Null.

Auf Grund der Ungleichungen (81) und des in § 13, 10. festgestellten Satzes über das Integral $J(\vartheta)$ gelangen wir nun durch Ausführung des Grenzprozesses (64) zu folgender *Rekursionsformel*:

$$(84) \quad \frac{n}{2} \frac{\gamma_n}{S_n} v_n = \frac{n}{n+1} v_{n-1}.$$

Da $v_1 = 1$ ist, erhalten wir unter Einführung des Wertes von γ_n (vgl. (54)) diese *Bestimmung von v_n* :

$$(85) \quad v_n = \frac{2}{n+1} \frac{\Gamma\left(\frac{2}{2}\right) \Gamma\left(\frac{3}{2}\right) \dots \Gamma\left(\frac{n}{2}\right)}{\left(\Gamma\left(\frac{1}{2}\right)\right)^{2+3+\dots+n}} S_2 S_3 \dots S_n.$$

Darin bedeutet Γ das Zeichen für die Gammafunktion und S_k steht für den Wert der unendlichen Reihe

$$1 + \frac{1}{2^k} + \frac{1}{3^k} + \frac{1}{4^k} + \dots.$$

§ 15. Maximale Dichte bei gitterförmiger Lagerung von Kugeln.

Von der Bestimmung des Volumens v_n machen wir eine *Anwendung auf die Frage der dichtesten Lagerung von Kugeln im n -dimensionalen Raume*.

Nehmen wir an, der größte Wert, den das Minimum $M(f)$ bei den sämtlichen positiven Formen f von der Determinante 1 überhaupt erreicht, sei ϱ_n , so enthält jede positive Formenklasse f von einer Determinante ≤ 1 wenigstens eine Form

$$\varphi(y_1, y_2, \dots, y_n) = \sum b_{hk} y_h y_k = b_{11} \left(y_1 + \frac{b_{12}}{b_{11}} y_2 + \dots + \frac{b_{1n}}{b_{11}} y_n \right)^2 + \psi(y_2, \dots, y_n),$$

wobei

$$0 < b_{11} \leq \varrho_n \sqrt[n]{D(f)}, \quad \pm \frac{b_{12}}{b_{11}} \leq \frac{1}{2}, \dots, \pm \frac{b_{1n}}{b_{11}} \leq \frac{1}{2}, \quad b_{12} \geq 0,$$

$$D(f) = b_{11} \text{Det.}(\psi) \leq 1$$

und ψ eine reduzierte Form der $n-1$ Variablen $y_2, \dots, y_n$ ist.

Das über den hierdurch definierten Bereich von Formen φ erstreckte $\frac{n(n+1)}{2}$ -fache Integral

$$\iint \dots \int db_{11} db_{12} \dots db_{nn}$$

wird infolgedessen mindestens so groß, ja, wie man leicht erkennt, größer als das Volumen des reduzierten Raumes der Formen mit Determinanten ≤ 1 sein. Dieses Integral hier findet sich

$$= \int_0^{\varrho_n} \frac{1}{2} b_{11}^{n-1} v_{n-1} \left(\left(\frac{1}{b_{11}} \right)^{\frac{n}{2}} - \left(\frac{b_{11}^{n-1}}{\varrho_n^n} \right)^{\frac{n}{2}} \right) db_{11} = \frac{v_{n-1}}{n+1} \varrho_n^{\frac{n}{2}}.$$

Daraus folgt

$$(86) \quad \frac{1}{2^n} \varrho_n^{\frac{n}{2}} \gamma_n > \frac{1}{2^{n-1}} S_n.$$

Wir können dieses Resultat folgendermaßen aussprechen:

Im n -dimensionalen Raume gibt es sicher solche parallelepipedische Lagerungen von lauter kongruenten Kugeln, daß der von den Kugeln erfüllte Raum mehr als das $\frac{1}{2^{n-1}} S_n$ -fache des ganzen unendlichen Raumes beträgt.

Dazu bemerken wir, daß bei einer „tetraedrischen“ Schichtung von Kugeln, welche der jedenfalls *extremen* Form

$$f = (x_1 + x_2 + \cdots + x_n)^2 + x_1^2 + x_2^2 + \cdots + x_n^2$$

entspricht, genau der

$$\frac{1}{2^{\frac{n}{2}} \sqrt{n+1}} \gamma_n = \frac{\pi^{\frac{n}{2}}}{2^{\frac{n}{2}} \sqrt{n+1} \Gamma\left(1 + \frac{n}{2}\right)} \text{-te Teil}$$

des unendlichen Raumes durch die Kugeln ausgefüllt wird, welcher Anteil *bei großem n wesentlich kleiner als der vorhin genannte Teil* ist, während allerdings in der Ebene und im Raume von 3 Dimensionen diese Schichtung nach regulären Dreiecken bzw. Tetraedern überhaupt die einzige dichteste gitterförmige Lagerung von kongruenten Kreisen bzw. Kugeln vorstellt.

§ 16. Asymptotisches Gesetz für die Klassenanzahlen ganzzahliger Formen.

Dem Volumen v_n kommt eine besondere *arithmetische Bedeutung* zu. Wir betrachten speziell die positiven Formen f mit lauter *ganzzahligen* Koeffizienten a_{hk} . Dabei ist immer $a_{11} \geq 1$. Mit Rücksicht hierauf haben wir den Ungleichungen (44), (45) zufolge zu einem gegebenen *positiven ganzzahligen Wert* $D(f) = D$ immer nur eine endliche Anzahl verschiedener reduzierter Formen mit ganzzahligen Koeffizienten und also auch *nur eine endliche Anzahl verschiedener Klassen ganzzahliger positiver Formen*. Diese *Klassenanzahl* für die Determinante D werde mit $H(D)$ bezeichnet.

Wir werden das asymptotische Gesetz nachweisen:

$$(87) \quad \lim_{D \rightarrow \infty} \left(\frac{H(1) + H(2) + \cdots + H(D)}{D^{\frac{n+1}{2}}} \right) = v_n.$$

1. Wir fassen in der Mannigfaltigkeit $\mathcal{A}$ aller quadratischen Formen $f = (a_{hk})$ diejenigen Punkte (a_{hk}^0) ins Auge, bei welchen *alle* $\frac{n(n+1)}{2}$

Koordinaten a_{hk}^0 ganze rationale Zahlen sind, und wir konstruieren um jeden solchen Punkt (a_{hk}) als Mittelpunkt den Würfelbereich

$$(88) \quad -\frac{1}{2} \leq a_{hk} - a_{hk}^0 \leq \frac{1}{2} \quad (h, k = 1, 2, \dots, n).$$

Wir erhalten ein *Netz von Würfeln*, welches die ganze Mannigfaltigkeit A einfach und lückenlos überdeckt.

Es sei nun D eine positive ganze Zahl, und wir bilden den Bereich $B(D, 1)$, der aus dem reduzierten Raume B durch die Forderungen

$$D(f) \leq D, \quad a_{11} \geq 1$$

herausgeschnitten wird. Es mögen N jener Würfel vollständig in diesen Bereich $B(D, 1)$ fallen und N^* darunter vorhanden sein, welche zwar nicht ganz in diesen Bereich fallen, aber doch in das Innere dieses Bereiches eintreten. *Zufolge der Ungleichungen (46) wird dann*

$$(89) \quad N < v_n D^{\frac{n+1}{2}}, \quad N + N^* > v_n D^{\frac{n+1}{2}} - \bar{v}_n D^{\frac{n}{2}}$$

sein.

Andererseits repräsentieren die Mittelpunkte der hier in der Anzahl N gezählten Würfel lauter *verschiedene Klassen ganzzahliger positiver Formen*, wobei die *Determinante* einen der Werte $1, 2, \dots, D$ hat, und wird jede Klasse solcher Formen durch *wenigstens einen* unter den Mittelpunkten jener $N + N^*$ Würfel repräsentiert. *Also ist*

$$(90) \quad N \leq H(1) + H(2) + \dots + H(D) \leq N + N^*.$$

Um das Resultat (87) zu gewinnen, wird nunmehr nach (89) und (90) eine solche Abschätzung der Größe N^* hinreichend sein, aus welcher sich

$$\lim_{D=\infty} \left(\frac{N^*}{D^{\frac{n+1}{2}}} \right) = 0$$

erschließen läßt.

2. Zu diesem Ende stellen wir zunächst einen einfach zu charakterisierenden Bereich fest, der alle jene $N + N^*$ Würfel vollständig in sich aufnimmt.

In einem beliebigen dieser Würfel gibt es stets solche Punkte (a_{hk}^*) , die in $B(D, 1)$ hineinfallen, für welche also insbesondere

$$(91) \quad 1 \leq a_{11}^* \leq a_{22}^* \leq \dots \leq a_{nn}^*, \quad \pm 2a_{hk}^* \leq a_{hh}^* \quad (h < k),$$

$$(92) \quad \lambda_n a_{11}^* a_{22}^* \dots a_{nn}^* \leq D$$

ist. Andererseits erfüllt im ganzen Bereiche des einzelnen Würfels jeder Punkt (a_{hk}) in bezug auf den *Mittelpunkt* (a_{hk}^0) des Würfels die Bedingungen (88). Namentlich gelten also diese Bedingungen für die Punkte $(a_{hk}) = (a_{hk}^*)$, und folgt daher

$$\frac{1}{2} \leq a_{hh}^* - \frac{1}{2} \leq a_{hh}^0.$$

Als ganze Zahlen sind hiernach die Werte a_{hh}^0 notwendig sämtlich ≥ 1 . Nunmehr haben wir im ganzen Würfel

$$\frac{1}{2} \leq a_{hh}^0 - \frac{1}{2} \leq a_{hh}.$$

Sodann folgt, immer mit Heranziehung von (91) und von (88),

$$a_{hh} \leq a_{hh}^* + 1 \leq a_{h+1, h+1}^* + 1 \leq a_{h+1, h+1} + 2 \leq 5 a_{h+1, h+1} \quad (h < n),$$

weiter

$$|a_{hk}| - 1 \leq |a_{hk}^*| \leq \frac{1}{2} a_{hh}^* \leq \frac{1}{2} (a_{hh} + 1), \quad |a_{hk}| \leq \frac{7}{2} a_{hh} \quad (h < k).$$

Endlich haben wir

$$a_{hh} \leq a_{hh}^* + 1 \leq 2 a_{hh}^*$$

und bringen diese Ungleichungen mit (92) in Verbindung.

Wir kommen damit zu folgendem Ergebnisse:

Die oben konstruierten $N + N^$ Würfel fallen mit allen ihren Punkten in den durch die Bedingungen*

$$(93) \quad \frac{1}{2} \leq a_{11}, \quad a_{hh} \leq 5 a_{h+1, h+1}, \quad |a_{hk}| \leq \frac{7}{2} a_{hh} \quad (h < k),$$

$$(94) \quad \frac{\lambda_n}{2^n} a_{11} a_{22} \dots a_{nn} \leq D$$

definierten Bereich hinein.

3. Der anschaulicheren Ausdrucksweise wegen wollen wir in der Mannigfaltigkeit A von der Größe einer Oberfläche und von orthogonalen Projektionen auf eine Ebene sprechen, indem wir der Definition dieser Begriffe die Formel

$$\sqrt{da_{11}^2 + da_{12}^2 + \dots + da_{nn}^2}$$

als Länge eines Linienelements zugrunde legen. Es wird sich nunmehr um die Oberfläche des durch (93), (94) definierten Körpers handeln, und wir werden für die Größe dieser Oberfläche eine obere Grenze von der Gestalt

$$(95) \quad \varpi_2 D \log D \text{ für } n = 2 \text{ bzw. } \varpi_n D^{\frac{n+1}{2} - \frac{1}{n}} \text{ für } n > 2$$

nachweisen, worin ϖ_n eine nur von n abhängende Konstante vorstellt. Im Falle $n = 2$ wollen wir uns hierbei $D \geq 2$ denken.

Im Volumenintegral des durch (93), (94) angewiesenen Körpers entspricht einem Intervall a_{11} bis $a_{11} + da_{11}$ des ersten Koeffizienten ein gewisser Beitrag

$$7^{n-1} a_{11}^{n-1} da_{11} \iint \dots \int da_{22} da_{23} \dots da_{nn}.$$

Projizieren wir andererseits diesen Körper orthogonal auf die Ebene $a_{11} = 0$ und beachten, wie diese Projektion sukzessive anwächst, während der Parameter a_{11} von $\frac{1}{2}$ an kontinuierlich zunimmt, so entspricht der

Zunahme von a_{11} bis $a_{11} + da_{11}$ des Parameters eine Vergrößerung jener Projektionsfläche um

$$d(7^{n-1}a_{11}^{n-1}) \int \int \dots \int da_{22} da_{23} \dots da_{nn},$$

wobei der Integrationsbereich von $a_{22}, a_{23}, \dots, a_{nn}$ der nämliche wie soeben ist. Dadurch zeigt sich, daß die durch das Gleichheitszeichen in (94) gelieferte begrenzende Fläche dieses Körpers auf die Ebene $a_{11} = 0$ eine Projektion ergibt, deren Flächeninhalt gewiß nicht größer ist als der Wert des Integrals

$$\int \int \dots \int (n-1) \frac{da_{11}}{a_{11}} da_{12} \dots da_{nn},$$

über den ganzen Körper erstreckt. *Hieraus resultiert zunächst für diese Projektion der Fläche (94) auf die Ebene $a_{11} = 0$ eine obere Grenze vom Typus (95).*

Nun hat in einem beliebigen Punkte (a_{hk}) der begrenzenden Fläche (94) die Tangentialebene an diese Fläche in laufenden Koordinaten (b_{hk}) die Gleichung

$$\frac{1}{n} \left(\frac{b_{11}}{a_{11}} + \frac{b_{22}}{a_{22}} + \dots + \frac{b_{nn}}{a_{nn}} \right) = 1;$$

wir ersehen daraus in Anbetracht der Ungleichungen (93), daß die Größe dieser Oberfläche ein gewisses, nur von n abhängendes Vielfaches ihrer Projektion auf die Ebene $a_{11} = 0$ nicht überschreitet.

Projizieren wir andererseits den ganzen durch (93), (94) definierten Körper auf die Ebene

$$(96) \quad b_{11} + b_{22} + \dots + b_{nn} = 0,$$

welche einer Tangentialebene der begrenzenden Fläche (94) parallel läuft, so wird die entstehende Projektion *einmal* genau von der Projektion dieser begrenzenden Fläche (94) geliefert und *ein zweites Mal* genau von den Projektionen aller übrigen durch die Ungleichungen (93) angewiesenen ebenen Seitenwände des Körpers überdeckt, wobei keine dieser ebenen Seitenwände orthogonal zur Ebene (96) ist. *Danach bestehen auch für die Flächeninhalte aller jener ebenen Seitenwände obere Grenzen vom Typus (95) und folgt endlich auch eine solche obere Grenze für die gesamte Oberfläche des in Rede stehenden Körpers.*

4. Wir kommen endlich auf die oben in der Anzahl N^* gezählten Würfel zurück. *Jeder dieser Würfel hat mit der Begrenzung des Bereichs $B(D, 1)$ Punkte gemein.* Diese Begrenzung besteht *erstens* aus einem Anteil der Fläche $D(f) = D$, *zweitens* aus Partien in den ebenen Seitenwänden des reduzierten Raumes, *drittens* aus einem Stück der Ebene $a_{11} = 1$.

Betrachten wir *zunächst diejenigen unter jenen N^* Würfeln, welche die Fläche $D(f) = D$ treffen*. Wir ordnen diese Würfel in *Serien* nach ihren Projektionen auf die Ebene $a_{11} = 0$. Es seien darunter zwei Würfel mit gleicher Projektion auf diese Ebene da, ihre Mittelpunkte seien $a_{11}^0, a_{12}^0, \dots, a_{nn}^0$ und $a_{11}^0 + d, a_{12}^0, \dots, a_{nn}^0$, so daß d eine positive ganze Zahl ist, und es seien $a_{11}, a_{12}, \dots, a_{nn}$ und $a_{11} + \varepsilon_{11}, a_{12} + \varepsilon_{12}, \dots, a_{nn} + \varepsilon_{nn}$ je ein Punkt in diesen Würfeln auf der Determinantenfläche $D(f) = D$. Dabei ist $d - 1 \leq \varepsilon_{11} \leq d + 1$ und sind $\varepsilon_{12}, \dots, \varepsilon_{nn}$ dem Betrage nach ≤ 1 ; für $h \geq 2$ folgt aus $a_{hh} + \varepsilon_{hh} \geq 1$ und $\varepsilon_{hh} \geq -1$ noch $a_{hh} + \varepsilon_{hh} \geq \frac{1}{2} a_{hh}$. Bilden wir nun die Differenz

$$\frac{1}{\partial \text{Det.} |a_{hk} + \varepsilon_{hk}|} (\text{Det.} |a_{hk} + \varepsilon_{hk}| - \text{Det.} |a_{hk}|) = 0,$$

$$\frac{\partial}{\partial a_{11}}$$

so hat darin ε_{11} den Koeffizienten 1; für die übrigen Glieder aber können wir vermöge (91) (vgl. auch (17) für $m = 1$) obere Grenzen angeben, die nur von n abhängen, so daß sich aus dieser Gleichung für $d + 1$ und damit auch *für die Maximalzahl der Würfel in einer jener Serien eine nur von n abhängende obere Grenze* ergibt. Das Volumen der in Rede stehenden Würfel übersteigt nun nicht das Produkt dieser Maximalzahl in die Projektion der Gesamtheit jener Würfel auf die Ebene $a_{11} = 0$, und letztere Projektion ist sicher nicht größer als die Projektion des ganzen, durch (93), (94) bestimmten Körpers auf die Ebene $a_{11} = 0$.

Nehmen wir *weiter diejenigen der N^* Würfel, welche eine bestimmte ebene Seitenwand*

$$(I, II) \quad \sum m_{hk} a_{hk} = 0$$

des reduzierten Raumes treffen. Wählen wir irgendeinen solchen Koeffizienten a_{hk} aus, für den der Faktor m_{hk} hier $\neq 0$ ist, so finden wir, daß die Anzahl derjenigen unter den betrachteten Würfeln, welche eine und die nämliche Projektion auf die Ebene $a_{hk} = 0$ liefern, nicht größer als eine gewisse aus den numerischen Faktoren in dieser Gleichung folgende Größe ist. Andererseits ist die gesamte Projektion aller dieser Würfel auf die Ebene $a_{hk} = 0$ nicht größer als die Projektion des ganzen Bereichs (93), (94) auf diese Ebene.

Endlich ist die Anzahl derjenigen unter den N^ Würfeln, welche die Ebene $a_{11} = 1$ durchsetzen, nicht größer als die Seitenwand $a_{11} = \frac{1}{2}$ des durch (93), (94) bestimmten Körpers.*

Aus allen diesen Umständen zusammen entnehmen wir für die Zahl N^* ebenfalls eine obere Grenze vom Typus (95), und wir gelangen dadurch zu folgendem Theoreme:

Die Gesamtanzahl aller verschiedenen Klassen ganzzahliger positiver quadratischer Formen mit n Variablen und von den Determinanten $1, 2, \dots, D$ ist zwischen zwei Grenzen

$$v_2 D^{\frac{s}{2}} \pm \omega_2 D \log D \text{ für } n = 2 \quad \text{bzw.} \quad v_n D^{\frac{n+1}{2}} \pm \omega_n D^{\frac{n+1}{2} - \frac{1}{n}} \text{ für } n > 2$$

eingeschlossen, wobei v_n das Volumen des Raumes der reduzierten Formen mit Determinanten ≤ 1 und ω_n eine gewisse andere positive, nur von n abhängende Konstante vorstellt.

Dabei ist im Falle $n = 2$ die Zahl $D \geq 2$ gedacht.

Inhalt.

	Seite
Problemstellung	53
§ 1. Charakter der positiven quadratischen Formen	55
§ 2. Anordnung in einer n -dimensionalen Mannigfaltigkeit	56
§ 3. Niedrigste Formen in einer Klasse	57
§ 4. Reduzierte Formen	59
§ 5. Die Wände des reduzierten Raumes	61
§ 6. Die Kanten des reduzierten Raumes	68
§ 7. Die Nachbarkammern des reduzierten Raumes	70
§ 8. Die Determinantenfläche	73
§ 9. Das Problem der dichtesten gitterförmigen Lagerung von Kugeln	74
§ 10. Bestimmung der extremen Formenklassen	76
§ 11. Die binären, ternären, quaternären Formen	78
§ 12. Volumen des reduzierten Raumes bis zur Determinantenfläche	80
§ 13. Verwendung Dirichletscher Reihen	82
§ 14. Auswertung des Volumens	90
§ 15. Maximale Dichte bei gitterförmiger Lagerung von Kugeln	94
§ 16. Asymptotisches Gesetz für die Klassenanzahlen ganzzahliger Formen	95

ZUR GEOMETRIE

Allgemeine Lehrsätze über die konvexen Polyeder.

(Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen.
Mathematisch-physikalische Klasse. 1897. S. 198—219.)

(Vorgelegt von David Hilbert in der Sitzung vom 31. Juli 1897.)

Ein *konvexer Körper* ist vollständig dadurch gekennzeichnet*), daß er eine abgeschlossene Punktmenge ist, innere Punkte besitzt, und daß jede gerade Linie, die innere Punkte von ihm aufnimmt, mit seiner Begrenzung stets zwei Punkte gemein hat (niemals mehr als zwei Punkte, falls auch die konvexen Körper, die sich ins Unendliche erstrecken, mit in Betracht gezogen werden). Infolge dieses einfachen Charakters spielen diese Gebilde eine gewisse Rolle bei der Behandlung einiger partieller Differentialgleichungen, die in der mathematischen Physik auftreten. Neuerdings habe ich in dem Buche „Geometrie der Zahlen“ gezeigt, daß auch merkwürdige arithmetische Beziehungen sich an die konvexen Körper knüpfen. Einen besonderen Reiz bieten die Sätze über konvexe Körper noch durch den Umstand dar, daß sie in der Regel für diese ganze Kategorie von Gebilden ohne jede Ausnahme Geltung haben.

Der vorliegende Aufsatz entstand bei Gelegenheit von Versuchen, den folgenden Satz zu beweisen, den ich seit längerer Zeit vermutete und dessen elementare Fassung nicht auf die Schwierigkeiten seiner Verifizierung schließen läßt: Wenn aus einer endlichen Anzahl von lauter *Körpern***) mit *Mittelpunkt*, die untereinander nur in den Begrenzungen zusammenstoßen, sich ein *konvexer Körper* aufbaut, so hat dieser stets ebenfalls einen *Mittelpunkt*.

Ich behandle hier nur diejenigen konvexen Körper, die ihre ganze Begrenzung in einer endlichen Anzahl von Ebenen liegen haben und auch

*) Geometrie der Zahlen, I. Lieferung, Leipzig bei B. G. Teubner, 1896; S. 200. Ich habe dort die betreffenden Gebilde *nirgends konkave* Körper genannt, hier will ich mich der kürzeren Bezeichnung *konvex* bedienen.

**) Unter den „Körpern mit Mittelpunkt“ dürfen hier, wie aus Lehrsatz V (S. 119) der Arbeit leicht ersichtlich ist, jedenfalls beliebige abgeschlossene Punktmenngen mit Mittelpunkt, denen eine bestimmte Größe der Oberfläche zukommt, verstanden werden.

sich nicht ins Unendliche erstrecken; ich entwickle über die eindeutige Festlegung eines derartigen *Polyeders* unter Verwendung der *Inhalte seiner Seitenflächen* einige Theoreme, die durch ihren leicht verständlichen Inhalt und andererseits die zu ihrem Nachweise erforderlichen Methoden Beachtung verdienen. Diese Theoreme sowie ihre Ausdehnung auf beliebige konvexe Körper werfen auch ein neues Licht auf die Eigenschaft der Kugel, unter allen Körpern von gleicher Oberfläche das größte Volumen zu besitzen.

§ 1. Vorbemerkungen.

1. Es seien rechtwinklige Koordinaten x, y, z zugrunde gelegt. Wenn von einer Richtung (α, β, γ) gesprochen wird, so soll gemeint sein, daß α, β, γ die Kosinus der Neigungswinkel der Richtung gegen die Richtungen der Koordinatenachsen sind. Es sei $\mathfrak{P}$ ein *konvexes Polyeder* mit n Seitenflächen, die in beliebiger Ordnung numeriert sein mögen. Es sei J das Volumen von $\mathfrak{P}$, $F_\nu (\nu = 1, \dots, n)$ der Flächeninhalt der ν^{ten} Seitenfläche, mithin $O = F_1 + \dots + F_n$ die Größe der Oberfläche von $\mathfrak{P}$. Es sei ferner $(\alpha_\nu, \beta_\nu, \gamma_\nu)$ die Richtung einer auf der ν^{ten} Seitenfläche nach dem *Äußeren* von $\mathfrak{P}$ hin errichteten Normalen. Die n Richtungen $(\alpha_\nu, \beta_\nu, \gamma_\nu)$ sind verschieden und jedenfalls derart, daß darunter sich irgend drei *unabhängige* finden, d. h. daß sie nicht alle einer einzigen Ebene angehören können.

Es sei p irgendein innerer Punkt von $\mathfrak{P}$, und es sei p_ν für $\nu = 1, \dots, n$ die Länge des von p auf die ν^{te} Seitenfläche von $\mathfrak{P}$ gefällten Perpendikels. Hält man den Punkt p und die Richtungen $(\alpha_\nu, \beta_\nu, \gamma_\nu)$ fest, betrachtet hingegen die Längen p_ν als veränderlich, so ändert sich das Polyeder $\mathfrak{P}$ und mit ihm sein Volumen J gemäß der Differentialformel:

$$(1) \quad dJ = F_1 dp_1 + \dots + F_n dp_n.$$

Will man diese Formel auf den Unterschied des Volumens derjenigen zwei Polyeder anwenden, die aus $\mathfrak{P}$ durch *Dilatation* vom Punkte p aus in allen Richtungen in einem Verhältnisse $t:1$, beziehlich $t+dt:1$ entstehen, so hat man F_ν durch $F_\nu t^2$ und dp_ν durch $p_\nu dt$ zu ersetzen. Wird die hervorgehende Formel nach t zwischen 0 und 1 integriert, so ergibt sich

$$(2) \quad 3J = F_1 p_1 + \dots + F_n p_n.$$

Benutzen wir statt p irgendeinen anderen Punkt in $\mathfrak{P}$, der die relativen Koordinaten a, b, c in bezug auf p haben mag, so haben wir p_ν durch $p_\nu - (a\alpha_\nu + b\beta_\nu + c\gamma_\nu)$ zu ersetzen. Weil a, b, c innerhalb gewisser Grenzen beliebig sind, so führt die Formel (2) zu

$$(3) \quad \sum F_\nu \alpha_\nu = 0, \quad \sum F_\nu \beta_\nu = 0, \quad \sum F_\nu \gamma_\nu = 0,$$

wo die Summen sich auf die Werte $\nu = 1, \dots, n$ beziehen. Die n Richtungen $(\alpha_\nu, \beta_\nu, \gamma_\nu)$ sind also weiter jedenfalls derart, daß diese drei Gleichungen (3) eine Auflösung in *positiven* Größen F_ν zulassen.

2. Es sei (α, β, γ) irgendeine Richtung und Θ das *Maximum*, ϑ das *Minimum* von $\varphi = \alpha x + \beta y + \gamma z$ im Bereiche des Polyeders $\mathfrak{P}$, so liegt $\mathfrak{P}$ zwischen den zwei parallelen Ebenen $\varphi = \vartheta$ und $\varphi = \Theta$ eingeschlossen und kann $\Theta - \vartheta = d$ als die *Breite* des Polyeders in der Richtung (α, β, γ) bezeichnet werden. Projiziert man die gesamte Oberfläche des Polyeders senkrecht auf die Ebene $\varphi = \vartheta$, so wird von der Projektion ein gewisses Polygon in dieser Ebene im ganzen Inneren doppelt überlagert, und ist danach der Flächeninhalt dieses Polygons gewiß $< \frac{1}{2} O$. Sodann schließt derjenige Zylinder, der senkrecht auf diesem Polygon steht und die Höhe d hat, so daß er von der Ebene $\varphi = \vartheta$ bis zur Ebene $\varphi = \Theta$ reicht, das Polyeder $\mathfrak{P}$ ganz in sich ein, und daraus folgt

$$(4) \quad \frac{1}{2} O d > J.$$

Es sei $\mathfrak{f}$ der *Schwerpunkt* von $\mathfrak{P}$ und s sein Abstand von der Ebene $\varphi = \Theta$. Nimmt man irgendeinen Punkt aus $\mathfrak{P}$ in der Ebene $\varphi = \vartheta$, so zerlegt sich das Polyeder $\mathfrak{P}$ in die Pyramiden, welche diesen Punkt als Spitze und die einzelnen, nicht durch ihn gehenden Seitenflächen von $\mathfrak{P}$ als Grundflächen haben; diese Pyramiden stoßen untereinander nur in den Begrenzungen zusammen. Da in einer Pyramide der Abstand des Schwerpunkts von der Basis $\frac{1}{4}$ der Höhe beträgt, so hat in jeder dieser Pyramiden der Schwerpunkt von der Ebene $\varphi = \Theta$ einen Abstand $\geq \frac{1}{4} d$, und daher ist auch $s \geq \frac{1}{4} d$; mit Hilfe von (4) folgt daher

$$(5) \quad s > \frac{J}{2O}.$$

Da dieses Resultat für jede beliebige Richtung (α, β, γ) gilt, so ist danach die Kugel vom Radius $\frac{J}{2O}$ mit $\mathfrak{f}$ als Mittelpunkt ganz im Inneren von $\mathfrak{P}$ enthalten, daraus folgt

$$J > \frac{4\pi}{3} \left(\frac{J}{2O} \right)^3, \quad \frac{O^2}{J^2} > \frac{\pi}{6}.$$

Betrachtet man weiter irgendeinen Punkt aus $\mathfrak{P}$ in der Ebene $\varphi = \vartheta$, irgendeinen Punkt aus $\mathfrak{P}$ in der Ebene $\varphi = \Theta$ und dazu den größten Kreis dieser Kugel in der parallelen Ebene $\varphi = \Theta - s$, so enthält das Polyeder sogleich die zwei Kegel, welche die Fläche dieses Kreises als Basis und ihre Spitze beziehlich in jenen zwei Punkten haben; daraus folgt

$$(6) \quad \frac{1}{3} \pi \left(\frac{J}{2O} \right)^2 d < J.$$

Ersetzt man (α, β, γ) durch $(-\alpha, -\beta, -\gamma)$, so vertauschen sich die Rollen der Ebenen $\varphi = \vartheta$ und $\varphi = \Theta$, und anstatt $s \geq \frac{1}{4}d$ gewinnt man die weitere Ungleichung $s \leq \frac{3}{4}d$.

Die Werte von s und d für die Richtung $(\alpha_v, \beta_v, \gamma_v)$ mögen s_v und d_v heißen. Für einen beliebigen Punkt p im Inneren von $\mathfrak{P}$ gilt offenbar stets

$$(7) \quad p_v < d_v.$$

Nimmt man endlich für den Punkt p , auf den sich (2) bezieht, den Schwerpunkt $\mathfrak{s}$ des Polyeders, so zeigt sich, daß der größte unter den Abständen s_v vorkommende Wert $\geq \frac{3J}{O}$ sein muß. Verbindet man diesen Umstand mit $s_v \leq \frac{3}{4}d_v$ und mit der Ungleichung (6), so folgt

$$(8) \quad \frac{O^3}{J^2} > \frac{\pi}{3}.$$

3. Als *konvexen Bereich* will ich überhaupt jede abgeschlossene Punktmenge bezeichnen, zu der mit irgend zwei Punkten stets auch die ganze sie verbindende Strecke gehört; ein *konvexer Körper* bedeutet dann einen solchen konvexen Bereich, der auch *innere* Punkte enthält, d. h. nicht ganz in einer Ebene gelegen ist.

Unter einer *Stützebene* eines konvexen Bereichs verstehe ich eine Ebene, die nicht zu beiden Seiten von sich Punkte des Bereichs liegen hat und selbst mindestens einen Punkt des Bereichs enthält. Ein konvexer Bereich besitzt durch jeden Punkt seiner Begrenzung wenigstens eine Stützebene. Wenn ferner ein konvexer Bereich sich *nicht ins Unendliche* erstreckt, gibt es zu jeder Richtung (α, β, γ) eine und nur eine Stützebene des Bereichs mit dieser Richtung als Normale und so, daß auf der Seite der Ebene, nach welcher die Richtung weist, kein Punkt des Bereichs liegt. Die Gleichung der betreffenden Stützebene ist

$$\alpha x + \beta y + \gamma z = r^*,$$

wenn r^* das *Maximum* von $\alpha x + \beta y + \gamma z$ im Bereiche bedeutet. —

Es seien jetzt irgend n verschiedene Richtungen $(\alpha_v, \beta_v, \gamma_v)$ für $v = 1, \dots, n$ gegeben, so, daß darunter drei unabhängige sich finden und daß die drei Gleichungen

$$(9) \quad \sum H_v \alpha_v = 0, \quad \sum H_v \beta_v = 0, \quad \sum H_v \gamma_v = 0$$

eine Lösung in *positiven* Werten H_v zulassen. Dann läßt sich zunächst zeigen, daß es jedenfalls ein Polyeder gibt mit n Seitenflächen, bei welchem jene Richtungen als die der äußeren Normalen der Flächen auftreten. In der Tat, durch die n Ungleichungen

$$(10) \quad \alpha_v x + \beta_v y + \gamma_v z - 1 \leq 0 \quad (v = 1, \dots, n)$$

wird ein konvexer Bereich Π definiert. Die n Ebenen $\alpha_v x + \beta_v y + \gamma_v z = 1$

sind Tangentialebenen der Kugel $x^2 + y^2 + z^2 \leq 1$. Also enthält Π diese Kugel, und es wird die Begrenzung des Bereichs Π von n , aber nicht schon von weniger Stützebenen geliefert. Ist jetzt x, y, z ein Punkt aus Π , so stellt $p_v = 1 - \alpha_v x - \beta_v y - \gamma_v z$ die Länge des von x, y, z auf die v^{te} jener Ebenen gefällten Perpendikels vor; unter Verwendung der vorausgesetzten Lösung von (9) folgt dann $\sum H_v p_v = \sum H_v$. Da die Größen H_v alle > 0 sind, ergibt sich hieraus, daß die Längen p_v nicht über eine gewisse Grenze hinausgehen. Da nun unter jenen n Ebenen sich drei solche finden, die sich nur in einem Punkte schneiden, ist danach Π in einem gewissen Parallelepipedum enthalten und kann sich also nicht ins Unendliche erstrecken; somit ist Π ein Polyeder, das der gestellten Forderung entspricht. Das Volumen von Π werde $= \frac{1}{6}$ gesetzt.

Wir bezeichnen die Linearform $\alpha_v x + \beta_v y + \gamma_v z$ mit φ_v . Es seien nun r_v für $v = 1, \dots, n$ irgendwelche n Größen ≥ 0 , und es sei etwa r der größte darunter vorkommende Wert; dann wird durch

$$(11) \quad \varphi_v - r_v \leq 0 \quad (v = 1, \dots, n)$$

ein konvexer Bereich — er heiße $\mathfrak{P}(r_v)$ — definiert, der ganz in demjenigen Polyeder liegt, welches durch Dilatation des Polyeders Π vom Nullpunkte aus im Verhältnisse $r : 1$ entsteht. Es kann $\mathfrak{P}(r_v)$ ein Polyeder mit n oder mit weniger Seitenflächen von nichtverschwindenden Inhalten werden oder auch sich auf die Fläche eines konvexen Polygons in einer Ebene oder auf eine Strecke oder gar auf den Nullpunkt allein reduzieren.

Wenn $\mathfrak{P}(r_v)$ nicht ein Polyeder mit n Polygonen als Begrenzung wird, so brauchen nicht alle n Ebenen $\varphi_v = r_v$ wirklich Punkte der Begrenzung des Polyeders zu enthalten. Es sei allgemein r_v^* das Maximum von φ_v im Bereiche $\mathfrak{P}(r_v)$, so ist jedenfalls $r_v^* \leq r_v$ und dann $\mathfrak{P}(r_v)$ identisch mit $\mathfrak{P}(r_v^*)$; die Ebenen $\varphi_v = r_v^*$ sind nunmehr sämtlich Stützebenen dieses Bereichs. Die so bestimmten Größen r_v^* mögen *tangentiale Parameter* von $\mathfrak{P}(r_v)$ heißen.

Ein solcher Bereich $\mathfrak{P}(r_v)$ wird nun, da er sich nicht ins Unendliche erstreckt, stets ein bestimmtes Volumen $J = J(r_v)$ besitzen, wobei jedenfalls

$$(12) \quad J \leq \left(\frac{r}{\varrho}\right)^3$$

sein wird; ferner wird das Gebiet von $\mathfrak{P}(r_v)$ in der Ebene $\varphi_v = r_v^*$, welches allgemein die v^{te} *Seitenfläche* von $\mathfrak{P}(r_v)$ genannt werden möge, einen bestimmten Flächeninhalt F_v besitzen; es kann J und jedes F_v auch Null sein. Diese Größen J und F_v ($v = 1, \dots, n$) sind offenbar im ganzen durch $r_1 \geq 0, \dots, r_n \geq 0$ definierten Gebiete *stetige* Funktionen von $r_1, \dots, r_n$.

Wenn für $\mathfrak{P}(r_v)$ alle Größen $F_v > 0$ ausfallen, sind die Werte r_v ohne

weiteres tangentielle Parameter. — Sind für zwei Bereiche $\mathfrak{P}(q_v)$ und $\mathfrak{P}(r_v)$ die Systeme q_v und r_v tangentielle Parameter, so stellen für $\mathfrak{P}((1-t)q_v + tr_v)$, wenn $0 < t < 1$ ist, die Werte $(1-t)q_v + tr_v$ ebenfalls tangentielle Parameter vor. Daraus ist zu erkennen, daß die Menge derjenigen Systeme r_v , welche tangentielle Parameter sind, einen konvexen Körper in der n -fachen Mannigfaltigkeit aller Systeme r_v bilden. Die Begrenzung dieses Körpers wird von einer endlichen Anzahl von Stützebenen geliefert, die sämtlich durch den Punkt $r_1 = 0, \dots, r_n = 0$ gehen, so daß der Körper ein Kegel mit diesem Punkte als Spitze ist; derselbe ist leicht mittels seiner Kanten zu charakterisieren, doch gehe ich auf diese Untersuchung, die für das Folgende entbehrlich ist, nicht weiter ein. — Wenn von zwei Bereichen $\mathfrak{P}(q_v)$ und $\mathfrak{P}(r_v)$, von denen keiner sich auf den Nullpunkt allein reduziere, der eine aus dem anderen durch *Dilatation und Translation* hervorgeht, d. h. beide *ähnlich und ähnlich gelegen* sind, so besteht zwischen ihren tangentialen Parametern q_v^* und r_v^* ein System von Gleichungen

$$(13) \quad q_v^* = a\alpha_v + b\beta_v + c\gamma_v + dr_v^*$$

mit bestimmten Werten a, b, c, d ; dabei ist noch stets $d > 0$; wenn $J(r_v) > 0$ ist, hat man $d = \frac{\sqrt[3]{J(q_v)}}{\sqrt[3]{J(r_v)}}$.

§ 2. Die Grundlagen der Untersuchung.

4. Herr Hermann Brunn*) hat den folgenden Satz entwickelt: Wenn ein konvexer Körper durch drei parallele Ebenen $\mathfrak{A}, \mathfrak{B}, \mathfrak{C}$ geschnitten wird, von denen die mittlere $\mathfrak{B}$ den Abstand zwischen $\mathfrak{A}$ und $\mathfrak{C}$ im Verhältnisse $t:1-t$ teilt, und es haben die Schnittfiguren des Körpers in $\mathfrak{A}, \mathfrak{B}, \mathfrak{C}$ die Flächeninhalte A, B, C , so besteht die Ungleichung

$$\sqrt{B} \geq (1-t)\sqrt{A} + t\sqrt{C};$$

dabei gilt hier das Zeichen $=$ nur dann, wenn der Teil des Körpers zwischen den Ebenen $\mathfrak{A}$ und $\mathfrak{C}$ sei es ein Zylinder, sei es ein Kegelstumpf mit den Grundflächen in diesen Ebenen, sei es ein Kegel mit der Grundfläche in der einen und der Spitze in der anderen dieser Ebenen ist. Eine entsprechende Eigenschaft der ebenen konvexen Figuren ist sehr einfach einzusehen, und die Methode von Brunn zum Nachweis jener Ungleichung ist wesentlich ein Schluß von 2 auf 3 Dimensionen, wobei die Schnitte des konvexen Körpers mit allen denjenigen Ebenen zu Hilfe genommen werden, welche die Schnittfigur in $\mathfrak{A}$ in einer Schar paralleler

*) *Über Ovale und Eiflächen*, S. 23, Inaugural-Dissertation, München 1887; *Über Kurven ohne Wendepunkte*, S. 50, Habilitationsschrift, München 1889.

Linien schneiden und gleichzeitig die Schnittfigur in $\mathfrak{C}$ jedesmal in zwei Stücke von gleichem Verhältnis der Flächeninhalte wie die Schnittfigur in $\mathfrak{A}$ zerlegen. Besondere Schwierigkeiten macht die strenge Erledigung der Grenzfälle, in welchen das Zeichen $=$ in jener Ungleichung eintritt.*)

Brunn hat auch bereits bemerkt, daß die eben erwähnten Sätze sich auf konvexe Körper in Mannigfaltigkeiten von mehr als drei Dimensionen ausdehnen lassen. Die hierzu erforderlichen Entwicklungen sind vollständig und in analytischer Form in den §§ 56—57 meiner „Geometrie der Zahlen“ auseinandergesetzt.

5. Hier nun werden uns die betreffenden Sätze für eine Mannigfaltigkeit von 4 Dimensionen dienlich sein; diese lassen sich auch leicht als Sätze über konvexe Körper in 3 Dimensionen fassen. Ich gehe wieder nur auf die Behandlung von Polyedern ein.

Es seien die n Richtungen $(\alpha_v, \beta_v, \gamma_v)$ für $v = 1, \dots, n$ wie in 3. beschaffen, und es sollen alle dort erklärten Bezeichnungen für sie Verwendung finden. Es seien q_v ($v = 1, \dots, n$) und r_v ($v = 1, \dots, n$) zwei Systeme von jedesmal n Größen ≥ 0 , so wird durch

$$0 \leq t \leq 1, \quad \alpha_v x + \beta_v y + \gamma_v z - (1-t)q_v - tr_v \leq 0 \quad (v = 1, \dots, n)$$

ein konvexer Bereich in der Mannigfaltigkeit der vier Variablen x, y, z, t definiert. Liegt dieser Bereich ganz in einer dreidimensionalen Ebene, so sind alle Größen $J((1-t)q_v + tr_v)$ für $0 \leq t \leq 1$ gleich Null. Anderenfalls haben wir, wenn wir den in Rede stehenden Satz auf die Schnitte dieses Bereichs mit den drei Ebenen anwenden, die durch $t = 0$, durch $t = 1$ und durch einen beliebigen Wert $t > 0$ und < 1 bestimmt sind, für letzteren Wert t die Ungleichung zu verzeichnen:

$$(14) \quad \sqrt[3]{J((1-t)q_v + tr_v)} \geq (1-t)\sqrt[3]{J(q_v)} + t\sqrt[3]{J(r_v)};$$

des weiteren tritt, wenn wir noch die q_v für $\mathfrak{P}(q_v)$ und die r_v für $\mathfrak{P}(r_v)$ als tangentielle Parameter voraussetzen, in dieser Ungleichung insbesondere das Gleichheitszeichen dann und nur dann ein, wenn alle q_v oder alle r_v Null sind oder die Bereiche $\mathfrak{P}(q_v)$ und $\mathfrak{P}(r_v)$ auseinander durch Dilatation und Translation hervorgehen, also Beziehungen

$$q_v = a\alpha_v + b\beta_v + c\gamma_v + dr_v \quad (d > 0)$$

statthaben. —

Sind $J(q_v)$ und $J(r_v)$ beide > 0 und wird $\frac{\sqrt[3]{J(q_v)}}{\sqrt[3]{J(r_v)}} = d$ gesetzt, so geht

*) Geometrie der Zahlen, §§ 56—57. — H. Brunn, Referat über eine Arbeit: Exakte Grundlagen für eine Theorie der Ovale, Sitzungsberichte der mathematisch-physikalischen Klasse der bayrischen Akademie der Wissenschaften, 1894, Bd. XXIV, S. 101.

(14) vermöge der Substitution $\frac{(1-t)d}{t} = \frac{1-\tau}{\tau}$ bei Multiplikation mit $\frac{\tau}{t}$ in

$$\sqrt[3]{J((1-\tau)\frac{q_v}{d} + \tau r_v)} \geq (1-\tau)\sqrt[3]{J\left(\frac{q_v}{d}\right)} + \tau\sqrt[3]{J(r_v)} = \sqrt[3]{J(r_v)}$$

über. Man erkennt daraus, daß die Ungleichung (14) wesentlich auf den einfacheren Satz hinausläuft: Hat man $J(q_v) = J(r_v)$, so gilt für $0 < t < 1$ stets $J((1-t)q_v + tr_v) \geq J(r_v)$.

6. Wir schreiben nun $\sqrt[3]{J((1-t)q_v + tr_v)} = j(t)$. Diese Funktion $j(t)$ ist im Intervalle $0 \leq t \leq 1$ eine stetige Funktion von t , und aus der allgemein aufgefaßten Regel (14) geht des weiteren

$$(15) \quad j(t) \geq \frac{t_1-t}{t_1-t_0} j(t_0) + \frac{t-t_0}{t_1-t_0} j(t_1)$$

für $0 \leq t_0 < t < t_1 \leq 1$ hervor. Diese Ungleichung lehrt, daß, wenn wir t, u als Parallelkoordinaten eines Punktes in einer zweidimensionalen Ebene $\mathfrak{E}$ deuten, durch $0 \leq t \leq 1, u = j(t)$, kurz ausgedrückt, ein nach der Seite der wachsenden u hin *konvexer Zug* in dieser Ebene geliefert wird; derselbe kann auch geradlinige Strecken aufweisen oder selbst eine einzige Strecke sein.

Nun wollen wir speziell annehmen, daß die q_v und die r_v und somit auch alle Systeme $(1-t)q_v + tr_v$ für $0 \leq t \leq 1$ tangentielle Parameter sind und weder alle q_v noch alle r_v Null sind, noch auch $\mathfrak{P}(q_v)$ und $\mathfrak{P}(r_v)$ auseinander durch Dilatation und Translation hervorgehen. Dann gilt nach den Ausführungen in 5. in der Ungleichung (15) stets das Zeichen $>$ und enthält daher der ebengenannte konvexe Zug keine geradlinige Strecke. Es sei J das Volumen, F_v der Flächeninhalt der v^{ten} Seitenfläche von $\mathfrak{P}((1-t)q_v + tr_v)$ und t dabei irgendein Wert > 0 und < 1 ; aus der Gleichung (1) entnimmt man dann leicht, daß durch

$$(1-\bar{t})(F_1 q_1 + \cdots + F_n q_n) + \bar{t}(F_1 r_1 + \cdots + F_n r_n) = 3J^{\frac{2}{3}}\bar{u},$$

$\bar{t}, \bar{u}$ als Koordinaten eines variablen Punktes in $\mathfrak{E}$ gedacht, die einzige vorhandene Tangente an diesen konvexen Zug im Punkte $t, u = j(t)$ dargestellt wird. Der Schnittpunkt dieser Tangente mit der Geraden $\bar{t} = 1$ hat die Ordinate

$$\bar{u} = \frac{F_1 r_1 + \cdots + F_n r_n}{3J^{\frac{2}{3}}};$$

nach der Natur eines nach der Seite der wachsenden u hin konvexen Zuges ohne geradlinige Strecken wird daher der Ausdruck

$$(16) \quad \frac{F_1 r_1 + \cdots + F_n r_n}{3J^{\frac{2}{3}}},$$

(in welchem $r_1, \dots, r_n$ fest und $J, F_1, \dots, F_n$ mit t variabel sind), eine

mit wachsendem t von $t = 0$ bis $t = 1$ beständig abnehmende Funktion von t sein. Insbesondere also wird dieser Ausdruck für $t = 0$ stets größer als für $t = 1$, d. h. $> J^{\frac{1}{3}}$ sein.

§ 3. Die einer Kugel umschriebenen Polyeder.

7. Nehmen wir speziell alle Größen $r_v = 1$, also für $\mathfrak{P}(r_v)$ das Polyeder Π aus 3., so geht der Ausdruck (16) in

$$\frac{O}{3 J^{\frac{2}{3}}}$$

über, wo J das Volumen, $O = F_1 + \dots + F_n$ die Gesamtoberfläche von $\mathfrak{P}((1-t)q_v + t)$ darstellt. Für das Polyeder Π ist zufolge (2): $3J = O$, also, wenn das Volumen von Π wie in der Zeile vorher mit $\frac{1}{\varrho^3}$ bezeichnet wird, $\frac{O}{3 J^{\frac{2}{3}}} = \frac{1}{\varrho}$.

Mit Bezug auf die Fälle, in welchen $\frac{O^3}{J^3}$ in der unbestimmten Form $\frac{0}{0}$ erscheint, sei folgendes bemerkt. Hat ein Bereich $\mathfrak{P}(p_v)$ ein Volumen $J > 0$ und sind die p_v tangentielle Parameter, so gilt nach (7) und (6) für ihn stets $p_v < \frac{12}{\pi} J^{\frac{1}{3}} \left(\frac{O^3}{J^3} \right)^{\frac{2}{3}}$. Läßt man nun die p_v als tangentielle Parameter sich stetig verändern und nach Grenzwerten konvergieren, die nicht sämtlich Null sind, während J dabei nach Null konvergiere, so wird daher das Verhältnis $\frac{O^3}{J^3}$ dabei stets über jede Grenze hinaus wachsen, selbst wenn auch O zugleich nach Null konvergiert.

Die hier erlangten Resultate sprechen wir folgendermaßen aus:

Lehrsatz I. Es seien $(\alpha_v, \beta_v, \gamma_v)$ für $v = 1, \dots, n$ irgend n verschiedene Richtungen so, daß darunter drei unabhängige sind und die Gleichungen

$$\sum H_v \alpha_v = 0, \quad \sum H_v \beta_v = 0, \quad \sum H_v \gamma_v = 0$$

eine Auflösung in n positiven Größen H_v zulassen. Dann und nur dann existieren Polyeder $\mathfrak{P}$ mit n oder weniger Seitenflächen, bei welchen die Richtungen der äußeren Normalen der Flächen zu jenen n Richtungen gehören. Unter diesen Polyedern gibt es solche mit n Flächen, die Kugeln umschrieben sind; es sind dies diejenigen, welche dem Polyeder $\alpha_v x + \beta_v y + \gamma_v z \leq 1$ ($v = 1, \dots, n$) ähnlich und ähnlich gelegen sind. Unter allen Polyedern $\mathfrak{P}$ haben diese letzteren das Minimum von $\frac{O^3}{J^2}$, des Verhältnisses der dritten Potenz der Oberfläche zum Quadrat des Volumens.

Ist ferner $\mathfrak{P}_0$ ein beliebiges unter den Polyedern $\mathfrak{P}$, das nicht zu den eben erwähnten speziellen Polyedern gehört, und konstruiert man zu den n Stützebenen von $\mathfrak{P}_0$ mit den Richtungen $(\alpha_v, \beta_v, \gamma_v)$ als äußeren Normalen Parallelebenen in einem Abstände l nach dem Äußeren des Polyeders hin, so begrenzen diese ein gewisses Polyeder $\mathfrak{P}_l$; dann ist für diese Polyeder $\mathfrak{P}_l$ die Funktion $\frac{O^3}{J^2}$ eine mit wachsendem l beständig abnehmende, und sie konvergiert für $l = \infty$ nach jenem Minimumwerte von $\frac{O^3}{J^2}$.

Dilatiert man vom Nullpunkte aus ein jedes Polyeder $\mathfrak{P}_l$ zu einem Polyeder $\bar{\mathfrak{P}}_l$ mit einer Oberfläche $= \frac{3}{q^3}$, so ist für diese Polyeder $\bar{\mathfrak{P}}_l$ das Volumen eine mit l beständig wachsende Größe, und es deckt sich $\bar{\mathfrak{P}}_\infty$ mit Π .

Der erste Teil des Lehrsatzes I ist bereits von Herrn L. Lindelöf*) durch interessante, ganz andersartige Betrachtungen bewiesen worden. Hier hat sich nicht allein die betreffende Maximeigenschaft des Polyeders Π herausgestellt, es hat sich zugleich für jedes Polyeder $\mathfrak{P}$, das nicht Π ähnlich und ähnlich gelegen ist, ein ganz bestimmter einfacher Übergang zu einem Polyeder dieser Art ergeben, wobei das Verhältnis $\frac{O^3}{J^2}$ beständig abnimmt und nach seinem Minimumwerte konvergiert.

Wenn man in derselben Weise, wie wir soeben die Ungleichung (14) und die Bemerkungen über das Eintreten des Gleichheitszeichens in ihr behandelt haben, von den entsprechenden Sätzen über beliebige konvexe Körper Gebrauch macht, so kommt man zu den Sätzen:

Unter allen konvexen Körpern besitzen die Kugeln das Minimum von $\frac{O^3}{J^2}$. Ist $\mathfrak{R}_0$ ein konvexer Körper, der keine Kugel vorstellt, und konstruiert man zu jeder Stützebene von $\mathfrak{R}_0$ eine parallele Ebene im Abstände l auf der dem Körper abgewandten Seite der Ebene, so begrenzen diese sämtlichen Parallelebenen jedesmal wieder einen konvexen Körper $\mathfrak{R}_l$; dann ist für diese Körper $\mathfrak{R}_l$ die Funktion $\frac{O^3}{J^2}$ eine mit wachsendem l beständig abnehmende, und sie konvergiert für $l = \infty$ nach ihrem Werte für Kugeln. Während für eine Kugel $O^3 = 36\pi J^2$ gilt, besteht danach für jeden konvexen Körper, der keine Kugel ist, die bekannte Ungleichung $O^3 > 36\pi J^2$.

Die Begrenzung von $\mathfrak{R}_l$ wird von der äußeren Parallelfläche im Abstände l zur Begrenzung von $\mathfrak{R}_0$ gebildet, und an diese Bemerkung knüpft sich leicht eine Ausdehnung der letzten Sätze auch auf nicht konvexe Körper, wie ich bei einer anderen Gelegenheit auseinandersetzen will.

*) Propriétés générales des polyèdres qui, sous une étendue superficielle donnée, renferment le plus grand volume, Mathematische Annalen, Bd. 2, S. 150. — Mémoire couronné par l'Académie Royale des Sciences de Berlin du prix Steiner en 1880.

§ 4. Bestimmung eines konvexen Polyeders unter Verwendung der Größen der Seitenflächen.

8. Es mögen alle Bezeichnungen wie in 5. Geltung haben, und man setze allgemein $\sqrt[3]{J(r_v)} = \psi(r_v)$; die Ungleichung (14) geht dann in

$$(17) \quad \psi((1-t)q_v + tr_v) \geq (1-t)\psi(q_v) + t\psi(r_v)$$

über. Es sei nun $\mathfrak{B}$ in der Mannigfaltigkeit der $n+1$ Variablen $r_1, \dots, r_n, w$ der durch

$$r_1 \geq 0, \dots, r_n \geq 0, \quad 0 \leq w \leq \psi(r_v)$$

definierte Bereich. Aus (17) ist zu ersehen, daß mit irgend zwei Punkten, die diesem Bereiche $\mathfrak{B}$ angehören, stets jeder Punkt der sie verbindenden Strecke zu ihm gehört. Da überdies wegen der Stetigkeit von $\psi(r_v)$ als Funktion der r_v dieser Bereich in jener Mannigfaltigkeit eine abgeschlossene Punktmenge ist, so stellt er einen *konvexen Körper* in derselben vor, freilich einen solchen, der sich auch ins Unendliche erstreckt. Wegen der Beziehung $\psi(tr_v) = t\psi(r_v)$ ($t \geq 0$) ist $\mathfrak{B}$ ein Kegel mit der Spitze im Nullpunkte $r_1 = 0, \dots, r_n = 0$.

Es sei weiter $\mathfrak{B}$ die Menge der durch

$$r_1 \geq 0, \dots, r_n \geq 0, \quad w = \psi(r_v)$$

definierten Punkte. Die *Begrenzung* von $\mathfrak{B}$ wird von den Punkten aus $\mathfrak{B}$, für welche $w = 0$ ist, und zudem von der Menge $\mathfrak{B}$ gebildet. Nach der Natur eines konvexen Körpers gibt es daher durch jeden Punkt von $\mathfrak{B}$ mindestens eine (n -dimensionale) *Stützebene* an $\mathfrak{B}$, also eine Ebene, die $\mathfrak{B}$ ganz auf einer Seite liegen hat, abgesehen von den Punkten aus $\mathfrak{B}$, die sie selbst enthält.

Sind $p_1, \dots, p_n$ lauter Werte > 0 , und ist J das Volumen, F der Flächeninhalt der v^{ten} Seitenfläche von $\mathfrak{P}(p_v)$, so besitzt, wie man leicht aus der Gleichung (1) erkennt, die Menge $\mathfrak{B}$ im Punkte $r_1 = p_1, \dots, r_n = p_n$, $w = \psi(p_v)$ die durch die Gleichung

$$3J^{\frac{2}{3}}w = F_1r_1 + \dots + F_nr_n$$

dargestellte Ebene als *Tangentialebene*. Diese Ebene ist somit die einzige Ebene durch den Punkt, welche überhaupt Stützebene an $\mathfrak{B}$ sein könnte, und demnach gilt dann für jeden beliebigen Punkt $r_1, \dots, r_n, w$ in $\mathfrak{B}$ stets

$$(18) \quad 3J^{\frac{2}{3}}w \leq F_1r_1 + \dots + F_nr_n.$$

9. Wir können nunmehr den folgenden Lehrsatz beweisen:

Lehrsatz II. *Es seien $(\alpha_v, \beta_v, \gamma_v)$ für $v = 1, \dots, n$ irgend n Richtungen, unter denen sich drei unabhängige finden, und F_v für $v = 1, \dots, n$ irgend n gegebene positive Größen so, daß*

$$\sum F_v \alpha_v = 0, \quad \sum F_v \beta_v = 0, \quad \sum F_v \gamma_v = 0$$

ist, endlich sei $\mathfrak{o}$ irgendein gegebener Punkt; dann existiert stets ein und nur ein konvexes Polyeder mit $\mathfrak{o}$ als Schwerpunkt und mit n Seitenflächen, wobei je eine Seitenfläche die Richtung $(\alpha_v, \beta_v, \gamma_v)$ als äußere Normale und F_v als Flächeninhalt hat.

Beweis. Wir setzen der Einfachheit halber $\mathfrak{o}$ als den Nullpunkt der Koordinaten voraus. Wir machen in bezug auf die gegebenen n Richtungen $(\alpha_v, \beta_v, \gamma_v)$ von den in 3. und 8. eingeführten Bezeichnungen Gebrauch und bilden dazu gemäß 8. die Punktmengen $\mathfrak{W}$ und $\mathfrak{B}$ in einer $n + 1$ -fachen Mannigfaltigkeit.

Zunächst wollen wir annehmen, daß ein Bereich $\mathfrak{P}(p_v)$ je mit F_v als Größe der v^{ten} Seitenfläche und mit $\mathfrak{o}$ als Schwerpunkt bereits bekannt ist, und wir beweisen, daß es nicht noch einen anderen Bereich derselben Art geben kann. Da alle $F_v > 0$ sein sollen, stellen die p_v gewiß *tangentiale Parameter* vor; da sie jedenfalls nicht alle Null sind, folgt aus (2): $J(p_v) > 0$, und da nunmehr der Schwerpunkt gewiß ein *innerer* Punkt in $\mathfrak{P}(p_v)$ ist, fallen die Größen p_v sämtlich > 0 aus. Nach (18) gilt für ein jedes System $r_1 \geq 0, \dots, r_n \geq 0$, da alsdann $r_1, \dots, r_n, w = \psi(r_v)$ ein Punkt in $\mathfrak{W}$ ist, stets

$$3(\psi(p_v))^3 \psi(r_v) \leq F_1 r_1 + \dots + F_n r_n.$$

Jetzt sei $\mathfrak{P}(q_v)$ gleichfalls ein Bereich mit $\mathfrak{o}$ als Schwerpunkt und F_v als Größe der v^{ten} Seitenfläche, so ist $F_1 q_1 + \dots + F_n q_n = 3(\psi(q_v))^3$ und folgt daher mit Rücksicht auf die vorstehende Ungleichung $\psi(p_v) \leq \psi(q_v)$. Genau so würde $\psi(q_v) \leq \psi(p_v)$ hervorgehen und also müßte zunächst $\psi(p_v) = \psi(q_v)$ sein. Dann würde also der Punkt $r_1 = q_1, \dots, r_n = q_n, w = \psi(q_v)$ in der Stützebene

$$3(\psi(p_v))^3 w = F_1 r_1 + \dots + F_n r_n$$

durch den Punkt $r_1 = p_1, \dots, r_n = p_n, w = \psi(p_v)$ an $\mathfrak{W}$ liegen. Mit diesen zwei Punkten in einer Stützebene müßte die ganze sie verbindende Strecke zur Begrenzung von $\mathfrak{W}$, also zu $\mathfrak{B}$, gehören; es würde demnach in der Ungleichung

$$\psi((1-t)p_v + tq_v) \geq (1-t)\psi(p_v) + t\psi(q_v) \quad \text{für } 0 < t < 1$$

stets das Gleichheitszeichen gelten. Dies hätte nach den Bemerkungen bei (14) zur Folge, daß das System q_v von der Form

$$q_v = a\alpha_v + b\beta_v + c\gamma_v + dp_v,$$

mit einem Koeffizienten $d > 0$ wäre. Dabei wäre nun d das Verhältnis $\sqrt[n]{J(q_v)} : \sqrt[n]{J(r_v)}$, also $= 1$, und da auch die Schwerpunkte von $\mathfrak{P}(p_v)$ und $\mathfrak{P}(q_v)$ übereinstimmen sollen, so hätte man weiter $a = 0, b = 0, c = 0$; also wäre $\mathfrak{P}(q_v)$ nicht von $\mathfrak{P}(p_v)$ verschieden.

Ich will jetzt den Schnitt von $\mathfrak{B}$ durch die Ebene $w=1$ mit $\mathfrak{B}'$ bezeichnen. Ferner bedeute $\frac{1}{\varrho^3}$ das Volumen des speziellen Polyeders $\mathfrak{P}(r_v=1)$; der Punkt $r_1=\varrho, \dots, r_n=\varrho, w=1$ liegt dann in $\mathfrak{B}'$ und $\mathfrak{B}$. Es mögen nun irgendwelche positive Werte F_v von der im Lehrsatz angegebenen Beschaffenheit vorausgesetzt werden; es sei F das Minimum unter diesen Werten. Für jeden Punkt $r_1=r'_1, \dots, r_n=r'_n, w=1$ in $\mathfrak{B}'$ gilt dann, wenn r' das Maximum unter den Werten r'_v bedeutet, mit Rücksicht auf (12):

$$(19) \quad F_1 r'_1 + \dots + F_n r'_n \geq F r' \geq F \varrho \sqrt[3]{J(r'_v)} \geq F \varrho.$$

Für den Punkt $r'_1=\varrho, \dots, r'_n=\varrho, w=1$ in $\mathfrak{B}'$ wird

$$F_1 r'_1 + \dots + F_n r'_n = (\sum F_v) \varrho.$$

Nun wird durch die Bedingung $F_1 r'_1 + \dots + F_n r'_n \leq (\sum F_v) \varrho$ aus $\mathfrak{B}'$ ein bestimmter konvexer Bereich ausgesondert, in dem für alle Koordinaten r'_v obere Grenzen bestehen. In diesem endlichen Bereich wird der Ausdruck $F_1 r'_1 + \dots + F_n r'_n$ ein bestimmtes Minimum besitzen, das zufolge (19) jedenfalls > 0 ausfallen wird und welches $3l^2$ heißen möge. Es sei $r_1=p'_1, \dots, r_n=p'_n, w=1$ ein Punkt aus $\mathfrak{B}'$, in dem dieses Minimum eintritt. Dieses Minimum ist zugleich das Minimum von $F_1 r'_1 + \dots + F_n r'_n$ im ganzen Bereiche $\mathfrak{B}'$, und also gilt in $\mathfrak{B}'$ stets $F_1 r'_1 + \dots + F_n r'_n \geq 3l^2$ und somit im Bereiche $\mathfrak{B}$, der ein Kegel mit der Spitze im Nullpunkte ist, stets $F_1 r_1 + \dots + F_n r_n \geq 3l^2 w$. Die Ebene

$$(20) \quad F_1 r_1 + \dots + F_n r_n = 3l^2 w$$

ist nunmehr eine Stützebene durch den Punkt $r_1=p'_1, \dots, r_n=p'_n, w=1$ an $\mathfrak{B}$, dieser Punkt somit jedenfalls ein Punkt aus $\mathfrak{B}$, mithin das Volumen von $\mathfrak{P}(p'_v)$ gleich 1. Es seien a, b, c die Koordinaten des Schwerpunktes von $\mathfrak{P}(p'_v)$ und allgemein

$$q'_v = p'_v - a\alpha_v - b\beta_v - c\gamma_v,$$

so sind wegen $J(p'_v)=1$ alle Größen $q'_v > 0$; es entsteht nun $\mathfrak{P}(q'_v)$ durch Translation aus $\mathfrak{P}(p'_v)$ und hat den Nullpunkt o als Schwerpunkt. Wegen der für die Größen F_v vorausgesetzten drei linearen Gleichungen liegt auch der Punkt $r_1=q'_1, \dots, r_n=q'_n, w=1$ in der Ebene (20). Da durch diesen Punkt nur eine Stützebene an $\mathfrak{B}$ geht, so leuchtet ein, daß für das Polyeder $\mathfrak{P}(q'_v)$ der Inhalt der v^{ten} Seitenfläche $= \frac{F_v}{l^2}$ ausfällt. Das Polyeder $\mathfrak{P}(lq'_v)$ ist dann ein solches mit F_v als Größe der v^{ten} Seitenfläche und o als Schwerpunkt, trägt mithin genau den im Lehrsatz geforderten Charakter.

10. Es seien die n Richtungen $(\alpha_v, \beta_v, \gamma_v)$ wieder so beschaffen, daß der Bereich $\alpha_v x + \beta_v y + \gamma_v z \leq 1$ sich nicht ins Unendliche erstreckt, und

es seien F_ν ($\nu = 1, \dots, n$) irgend n Größen ≥ 0 , so daß $\sum F_\nu \alpha_\nu = 0$, $\sum F_\nu \beta_\nu = 0$, $\sum F_\nu \gamma_\nu = 0$ ist. Es sollen diese Größen nicht sämtlich Null sein, so daß $F_1 + \dots + F_n = O > 0$ ist; sie brauchen aber jetzt nicht sämtlich > 0 zu sein.

Wir betrachten diejenigen Richtungen $(\alpha_\nu, \beta_\nu, \gamma_\nu)$, zu denen ein $F_\nu > 0$ gegeben ist. Haben wir *erstens* den Fall, daß unter diesen Richtungen schon drei unabhängige vorkommen, so gibt es nach dem Lehrsatz II unter den Bereichen $\mathfrak{P}(r_\nu)$ zu den gegebenen n Richtungen ein und nur ein Polyeder $\mathfrak{P}(p_\nu)$ je mit F_ν als Größe der ν^{ten} Seitenfläche für $\nu = 1, \dots, n$ und noch mit beliebigem Schwerpunkte; wir wollen dann unter $J(F_\nu)$ das Volumen dieses Polyeders verstehen. *Zweitens* mögen dagegen alle jene Richtungen $(\alpha_\nu, \beta_\nu, \gamma_\nu)$, für welche ein $F_\nu > 0$ gegeben ist, einer einzigen Ebene E angehören. Nähern wir uns dann dem gegebenen Systeme F_ν irgendwie vermittlest solcher Systeme $F_\nu^{(0)}$, die dem zuerst genannten Falle entsprechen, und konstruieren für diese jedesmal das zugehörige Polyeder $\mathfrak{P}(p_\nu^{(0)})$ wie soeben, so konvergiert für diese Polyeder $\mathfrak{P}(p_\nu^{(0)})$ die senkrechte Projektion ihrer Oberfläche auf die Ebene E schließlich nach Null; es wird damit für diese Polyeder auch die kleinste unter ihren *Breiten* d (s. 2.) in den Richtungen dieser Ebene und zufolge der Formel (4): $\frac{1}{2} Od > J$ also auch $J(F_\nu^{(0)})$ stets nach Null konvergieren. In diesem zweiten Falle setzen wir demgemäß $J(F_\nu) = 0$. Endlich werde, wenn alle Größen $F_\nu = 0$ sind, ebenfalls $J(F_\nu) = 0$ gesetzt.

Auf solche Weise ist nun für jedes System F_ν in dem durch

$$(21) \quad F_\nu \geq 0, \quad \sum F_\nu \alpha_\nu = 0, \quad \sum F_\nu \beta_\nu = 0, \quad \sum F_\nu \gamma_\nu = 0$$

definierten Bereiche der Wert $J(F_\nu)$ eindeutig festgelegt und stellt dieser Wert, wie aus dem Lehrsatz II und den eben gemachten Bemerkungen leicht ersichtlich ist, eine stetige Funktion der F_ν in diesem ganzen Bereiche vor.

Wir setzen nun $(J(F_\nu))^{\frac{2}{3}} = \Psi(F_\nu)$. Dann gilt, wenn G_ν und H_ν ($\nu = 1, \dots, n$) irgend zwei Systeme in dem Bereiche (21) sind und noch $\Psi(H_\nu) > 0$ ist, für jeden Wert $t > 0$ und < 1 stets

$$\Psi((1-t)G_\nu + tH_\nu) \geq (1-t)\Psi(G_\nu) + t\Psi(H_\nu),$$

und zwar tritt das Zeichen $=$ hier nur dann ein, wenn $G_1 : \dots : G_n = H_1 : \dots : H_n$ ist.

Daß in dem zuletzt bezeichneten Falle diese Ungleichung und zwar mit dem Zeichen $=$ erfüllt ist, leuchtet ohne weiteres ein. Nehmen wir nun an, es sei nicht $G_1 : \dots : G_n = H_1 : \dots : H_n$. Nach (18) gilt für jedes System von Größen $r_\nu \geq 0$ stets

$$(22) \quad G_1 r_1 + \dots + G_n r_n \geq 3\Psi(G_\nu)\psi(r_\nu),$$

$$(23) \quad H_1 r_1 + \dots + H_n r_n \geq 3\Psi(H_\nu)\psi(r_\nu).$$

Wegen $\Psi(H_\nu) > 0$ und $t > 0$ gibt es ein bestimmtes Polyeder $\mathfrak{P}(p_\nu)$ mit $(1-t)G_\nu + tH_\nu$ als Größe der ν^{ten} Seitenfläche und dem Nullpunkt als Schwerpunkt. Für dieses Polyeder hat man dann

$$((1-t)G_1 + tH_1)p_1 + \dots + ((1-t)G_n + tH_n)p_n = 3\Psi((1-t)G_\nu + tH_\nu)\psi(p_\nu).$$

Es sind dabei die p_ν sämtlich > 0 und geht daher durch den Punkt $r_1 = p_1, \dots, r_n = p_n$, $w = \psi(p_\nu)$ nur eine Stützebene an $\mathfrak{B}$; nun gelten die Ungleichungen (22), (23) auch für $r_1 = p_1, \dots, r_n = p_n$; aus dem eben angeführten Grunde und weil nicht

$$(1-t)G_1 + tH_1 : \dots : (1-t)G_n + tH_n = H_1 : \dots : H_n$$

ist, hat dabei in der zweiten jedenfalls das Zeichen $>$ statt. Man erhält somit aus ihnen

$$((1-t)G_1 + tH_1)p_1 + \dots + ((1-t)G_n + tH_n)p_n > 3((1-t)\Psi(G_\nu) + t\Psi(H_\nu))\psi(p_\nu);$$

der Vergleich dieser Relation mit der davor angegebenen liefert unmittelbar die zu beweisende Ungleichung.

Es sei jetzt O eine beliebige positive Größe. Unter allen Polyedern $\mathfrak{P}(r_\nu)$ mit einer Gesamtoberfläche $= O$ gibt es, wie schon in 7. ausgeführt wurde, ein, bis auf Translationen völlig bestimmtes Polyeder wirklich mit n Seitenflächen, welches einer Kugel umschrieben ist. Es sei Φ_ν die Größe der ν^{ten} Seitenfläche bei diesem Polyeder. Ist dann F_ν ($\nu = 1, \dots, n$) irgendein von dem Systeme der Φ_ν ($\nu = 1, \dots, n$) verschiedenes System von Größen ≥ 0 im Bereiche (21) und gleichfalls mit der Summe $F_1 + \dots + F_n = O$, so gilt nach dem Lehrsatz I stets $\Psi(F_\nu) < \Psi(\Phi_\nu)$. Betrachtet man nun t, u als Parallelkoordinaten in einer Ebene und faßt die Punkte $0 \leq t \leq 1, u = \Psi((1-t)F_\nu + t\Phi_\nu)$ ins Auge, so bilden diese Punkte nach den vorhin gewonnenen Ungleichungen daselbst einen nach der Seite der wachsenden u hin konvexen Zug, und nach der eben gemachten Bemerkung hat dabei u für $t = 1$ seinen größten Wert. Nach der Natur eines solchen Zuges muß nun, wenn auf demselben u zugleich mit t am größten ist, auf seiner ganzen Ausdehnung u mit abnehmendem t beständig abnehmen. Danach stellt $J((1-t)F_\nu + t\Phi_\nu)$ im Intervalle $0 \leq t \leq 1$ eine mit wachsendem t beständig wachsende Funktion vor. Damit ist ein sehr bemerkenswerter neuer Prozeß gefunden, um von einem beliebigen der Polyeder $\mathfrak{P}(r_\nu)$, welches nicht einer Kugel und zwar mit n Berührungen umschrieben ist, zu einem Polyeder $\mathfrak{P}(r_\nu)$ dieser besonderen Art überzugehen so, daß die Oberfläche sich nicht ändert und das Volumen beständig wächst.

§ 5. Konvexe Körper mit Mittelpunkt.

11. Es sei jetzt n eine gerade Zahl $= 2m$ und die $2m$ Richtungen $(\alpha_\nu, \beta_\nu, \gamma_\nu)$ ($\nu = 1, \dots, 2m$) sollen aus m Paaren entgegengesetzter Rich-

tungen bestehen. Sowie sich unter diesen $2m$ Richtungen drei unabhängige finden, was wir jetzt voraussetzen wollen, zeigt sich bereits, daß der Bereich $\alpha_\nu x + \beta_\nu y + \gamma_\nu z - 1 \leq 0$ ($\nu = 1, \dots, n$) ganz im Endlichen liegt; denn es begrenzen alsdann die sechs Ebenen $\alpha_\nu x + \beta_\nu y + \gamma_\nu z = 1$ zu diesen drei Richtungen und den drei ihnen entgegengesetzten ein Parallelepipedum, welches jenen Bereich ganz in sich schließt.

Es sei

$$(24) \quad \alpha_{m+\mu} = -\alpha_\mu, \quad \beta_{m+\mu} = -\beta_\mu, \quad \gamma_{m+\mu} = -\gamma_\mu \quad (\mu = 1, \dots, m).$$

Wir wollen nun von den Bereichen $\mathfrak{P}(r_\nu)$ zu den $2m$ gegebenen Richtungen nur diejenigen betrachten, bei welchen $r_{m+\mu} = r_\mu$ für $\mu = 1, \dots, m$ ist; einen solchen Bereich bezeichnen wir durch $\mathfrak{P}\{r_\mu\}$, er ist stets ein Bereich mit dem Nullpunkt als *Mittelpunkt*, und haben wir dabei stets $F_{m+\mu} = F_\mu$ ($\mu = 1, \dots, m$), unter F_ν die Größe der ν^{ten} Seitenfläche des Bereichs verstanden. Bezeichnen wir das Volumen von $\mathfrak{P}\{r_\mu\}$ mit $J\{r_\mu\}$, so folgt aus (14) sogleich

$$\sqrt[3]{J\{(1-t)q_\mu + tr_\mu\}} \geq (1-t)\sqrt[3]{J\{q_\mu\}} + t\sqrt[3]{J\{r_\mu\}},$$

und auch die Bemerkungen über das Eintreten des Gleichheitszeichens in (14) sind sinngemäß auf diese Ungleichung zu übertragen. Setzen wir $\sqrt[3]{J\{r_\mu\}} = \psi\{r_\mu\}$, so ist danach der durch

$$r_1 \geq 0, \dots, r_m \geq 0, \quad 0 \leq w \leq \psi\{r_\mu\}$$

definierte Bereich in der Mannigfaltigkeit der $m+1$ Variablen $r_1, \dots, r_m, w$ ein *konvexer Körper*; und durch jeden Punkt, wo $r_1 > 0, \dots, r_m > 0, w = \psi\{r_\mu\}$ ist, gibt es stets nur eine Stützebene an diesen Körper.

Erwägen wir nun, daß für beliebige $2m$ Größen $F_\nu \geq 0$ ($\nu = 1, \dots, 2m$), bei welchen $F_{m+\mu} = F_\mu$ ($\mu = 1, \dots, m$) ist, wegen (24) die Gleichungen $\sum F_\nu \alpha_\nu = 0, \sum F_\nu \beta_\nu = 0, \sum F_\nu \gamma_\nu = 0$ stets erfüllt sind, so gelangen wir durch ganz entsprechende Überlegungen wie in 9. zu dem Satze:

Lehrsatz III. *Es seien $(\alpha_\mu, \beta_\mu, \gamma_\mu)$ für $\mu = 1, \dots, m$ irgend m verschiedene Richtungen, von denen auch keine zwei einander entgegengesetzt sind und unter denen sich drei unabhängige finden, ferner seien F_μ für $\mu = 1, \dots, m$ irgend m positive Größen, und $\mathfrak{o}$ ein gegebener Punkt; dann gibt es stets ein und nur ein konvexes Polyeder mit $\mathfrak{o}$ als Mittelpunkt und mit $2m$ paarweise parallelen Seitenflächen, von denen je ein Paar als Richtungen der äußeren Normalen $(\alpha_\mu, \beta_\mu, \gamma_\mu)$ und $(-\alpha_\mu, -\beta_\mu, -\gamma_\mu)$ und als Größe der Seitenfläche F_μ haben.*

Wir ziehen hieraus und aus Lehrsatz II sogleich die weitere Folgerung:

Lehrsatz IV. *Ein konvexes Polyeder mit einer geraden Anzahl von Seitenflächen, wobei diese paarweise parallel und von gleichem Flächeninhalt sind, ist stets ein Polyeder mit Mittelpunkt.*

Denn es sei $n = 2m$ die Anzahl der Seitenflächen des Polyeders, $(\alpha_\nu, \beta_\nu, \gamma_\nu)$ für $\nu = 1, \dots, n$ die Richtung der äußeren Normalen, F_ν die Größe seiner ν^{ten} Seitenfläche, und man habe $\alpha_{m+\mu} = -\alpha_\mu$, $\beta_{m+\mu} = -\beta_\mu$, $\gamma_{m+\mu} = -\gamma_\mu$, $F_{m+\mu} = F_\mu$ ($\mu = 1, \dots, m$), endlich sei o der Schwerpunkt des Polyeders. Nach dem Lehrsatz II kann es überhaupt nur ein konvexes Polyeder, also nur das vorgelegte geben, bei welchem alle die eben erwähnten Stücke in der betreffenden Weise eintreten; andererseits ist nach Lehrsatz III zu diesen Stücken speziell ein konvexes Polyeder mit o als Mittelpunkt vorhanden; mithin ist das vorgelegte Polyeder notwendig ein Polyeder mit Mittelpunkt.

12. Wir können weiter den Satz aufstellen:

Lehrsatz V. Wenn irgendwelche (nicht notwendig konvexe) Polyeder in endlicher Anzahl, von denen jedes einen Mittelpunkt hat und die untereinander nur in Punkten der Begrenzungen zusammenstoßen, durch ihre Vereinigung ein konvexes Polyeder erfüllen, so hat dieses zusammengesetzte konvexe Polyeder stets ebenfalls einen Mittelpunkt.

Denn betrachten wir irgendeine Seitenfläche $\mathfrak{F}$ dieses so zusammengesetzten konvexen Polyeders $\mathfrak{P}$. Es sei (α, β, γ) die Richtung der äußeren Normale von $\mathfrak{F}$. Unter den Einzelpolyedern, deren Vereinigung $\mathfrak{P}$ vorstellt, finden sich dann notwendig ebenfalls solche, welche sei es eine, sei es mehrere Seitenflächen mit (α, β, γ) als Richtung der äußeren Normalen darbieten. Bei jedem hier in Betracht kommenden Einzelpolyeder treten, da das Polyeder jedesmal einen Mittelpunkt besitzt, symmetrisch in bezug auf diesen, zu den Seitenflächen mit der äußeren Normalenrichtung (α, β, γ) ebensoviele Seitenflächen mit der äußeren Normalenrichtung $(-\alpha, -\beta, -\gamma)$ auf; und irgend zwei einander auf diese Weise entsprechende Seitenflächen haben stets gleichen Flächeninhalt. Bilden wir den gesamten Flächeninhalt aller bei den Einzelpolyedern auftretenden Seitenflächen mit der äußeren Normalenrichtung (α, β, γ) und subtrahieren davon den gesamten Flächeninhalt aller bei ihnen auftretenden Seitenflächen mit der äußeren Normalenrichtung $(-\alpha, -\beta, -\gamma)$, so muß daher die Differenz $= 0$ sein. Nun wird, soweit diese verschiedenen Seitenflächen im Inneren von $\mathfrak{P}$ liegen, hier die Gesamtheit der Seitenflächen der ersteren Stellung genau überdeckt von der Gesamtheit der Seitenflächen der anderen Stellung; also verschwindet für sich der Teil jener Differenz, welcher sich auf Seitenflächen bezieht, die (abgesehen vielleicht von Punkten ihres Randes) ins Innere von $\mathfrak{P}$ fallen. Weiter setzen diejenigen von den Seitenflächen der ersteren Stellung, welche auf die Begrenzung von $\mathfrak{P}$ fallen, hier eben die Seitenfläche $\mathfrak{F}$ von $\mathfrak{P}$ zusammen. Nunmehr leuchtet ein, daß noch Seitenflächen der anderen Stellung übrig bleiben, welche zusammen eine begrenzende Seitenfläche

von $\mathfrak{P}$ mit der äußeren Normalenrichtung $(-\alpha, -\beta, -\gamma)$ und genau von einem Flächeninhalt gleich dem von $\mathfrak{F}$ ergeben müssen. Es sind danach die Seitenflächen des konvexen Polyeders $\mathfrak{P}$ paarweise parallel und von gleichem Flächeninhalt. Nach dem Lehrsatz IV ist somit $\mathfrak{P}$ ein Polyeder mit Mittelpunkt.

§ 6. Konvexe Restbereiche.

Die Gesamtheit der Punkte x, y, z , für welche sowohl x , wie y , wie z ganze rationale Zahlen sind, soll das *Zahlengitter* heißen; ein einzelner Punkt daraus heiße ein *Gitterpunkt*. Unter einem *konvexen Restbereich* soll ein konvexer Körper $\mathfrak{R}$ von solcher Art verstanden werden, daß $\mathfrak{R} = \mathfrak{R}_{0,0,0}$ und die Gesamtheit derjenigen Körper $\mathfrak{R}_{a,b,c}$, die aus $\mathfrak{R}_{0,0,0}$ durch die Translationen vom Nullpunkte $0, 0, 0$ nach den verschiedenen anderen Gitterpunkten a, b, c hervorgehen, den ganzen Raum *lückenlos* überdecken, und zwar so, daß irgend zwei von diesen Körpern *höchstens in Punkten der Begrenzung* zusammenstoßen.

Lehrsatz VI. *Ein jeder konvexe Restbereich ist ein konvexes Polyeder mit Mittelpunkt und wird von nicht mehr als $2(2^3 - 1)$ Seitenflächen begrenzt; dabei ist weiter jede Seitenfläche ein konvexes Polygon mit Mittelpunkt.*

Beweis. Es sei $\mathfrak{R} = \mathfrak{R}_{0,0,0}$ ein konvexer Restbereich; man erkennt sofort, daß $\mathfrak{R}$ nicht einen ganz im Endlichen gelegenen konvexen Körper von einem Volumen > 1 enthalten kann und somit selbst ganz im Endlichen liegen muß; also wird $\mathfrak{R}$ auch nur mit einer endlichen Anzahl der anderen Körper $\mathfrak{R}_{a,b,c}$ in Punkten der Begrenzung zusammenstoßen. Da der Körper $\mathfrak{R}$ von jedem dieser Körper $\mathfrak{R}_{a,b,c}$, mit dem er zusammentrifft, durch eine gemeinsame Stützebene geschieden werden kann, so daß die gemeinschaftlichen Punkte beider Körper in dieser Ebene und im übrigen der eine ganz auf der einen, der andere ganz auf der anderen Seite von ihr liegt, so leuchtet zuvörderst ein, daß $\mathfrak{R}$ jedenfalls ein von einer endlichen Anzahl von Ebenen begrenztes konvexes Polyeder ist.

Wir wollen unter $\mathfrak{L}$ denjenigen Körper verstehen, der zu $\mathfrak{R}$ symmetrisch in bezug auf den Nullpunkt liegt; dann ist $\mathfrak{L}$ ebenfalls ein konvexer Restbereich, so daß die sämtlichen Körper $\mathfrak{L}_{a,b,c}$, die aus $\mathfrak{L}$ durch die Translationen nach den einzelnen Gitterpunkten a, b, c entstehen, den ganzen Raum erfüllen und dabei je zwei unter ihnen stets in den inneren Punkten durchweg verschieden sind. Es wird nun unter allen diesen Polyedern $\mathfrak{L}_{a,b,c}$ eine endliche Anzahl von solchen geben, welche ins Innere von $\mathfrak{R}$ eintreten; und der Körper $\mathfrak{R}$ erscheint dann genau zusammengesetzt aus den einzelnen Polyedern, welche $\mathfrak{R}$ mit diesen einzelnen Körpern $\mathfrak{L}_{a,b,c}$ gemein hat. Nun ist ein Bereich $\mathfrak{L}_{a,b,c}$ jedesmal symmetrisch zu $\mathfrak{R}$ in bezug auf den Punkt $\frac{a}{2}, \frac{b}{2}, \frac{c}{2}$; ein Polyeder, welches

$\mathfrak{R}$ und $\mathfrak{L}_{a,b,c}$ gemein haben, wird danach ein Polyeder mit $\frac{a}{2}, \frac{b}{2}, \frac{c}{2}$ als Mittelpunkt sein. Es erscheint also $\mathfrak{R}$ zerlegt in eine endliche Anzahl von Polyedern mit Mittelpunkt; nach dem Lehrsatz V ist daher $\mathfrak{R}$ selbst ein Polyeder mit Mittelpunkt.

Da durch eine Translation eines konvexen Restbereichs offenbar stets wieder ein solcher Bereich hervorgeht, so wollen wir jetzt der Einfachheit wegen annehmen, es habe $\mathfrak{R}$ den Nullpunkt als Mittelpunkt. Betrachten wir nun irgendeine Seitenfläche $\mathfrak{S}$ von $\mathfrak{R} = \mathfrak{R}_{0,0,0}$, so gibt es unter allen übrigen Polyedern $\mathfrak{R}_{a,b,c}$ eines oder mehrere, welche sich an diese Seitenfläche mit einem Flächenstück (nicht bloß mit Punkten einer Kante) anlegen. Ist $\mathfrak{R}_{a,b,c}$ ein derartiges Polyeder, so ist das Flächenstück aus $\mathfrak{S}$, das $\mathfrak{R}$ und $\mathfrak{R}_{a,b,c}$ gemein haben, da $\mathfrak{R}_{a,b,c}$ symmetrisch zu $\mathfrak{R}$ in bezug auf den Punkt $\frac{a}{2}, \frac{b}{2}, \frac{c}{2}$ ist, ein Polygon mit $\frac{a}{2}, \frac{b}{2}, \frac{c}{2}$ als Mittelpunkt. Danach erscheint das konvexe Polygon $\mathfrak{S}$ zerlegt in Polygone mit Mittelpunkt, und von diesen ist noch leicht ersichtlich, daß sie untereinander nur in den Rändern zusammentreffen können. Nun gilt ein dem Satze V ganz entsprechender Satz für zwei Dimensionen, und danach ist die Fläche $\mathfrak{S}$ notwendig selbst ein Polygon mit Mittelpunkt.

Wie sich ferner ergeben hat, liegt auf der Fläche $\mathfrak{S}$, noch von ihrem Rande abgesehen, mindestens ein Punkt $\frac{a}{2}, \frac{b}{2}, \frac{c}{2}$, wo a, b, c ganze Zahlen sind. Dabei können a, b, c niemals sämtlich gerade Zahlen sein, weil die Gitterpunkte im Inneren der betrachteten Polyeder liegen.

Andererseits kann kein Punkt $\frac{a}{2}, \frac{b}{2}, \frac{c}{2}$, bei dem a, b, c ganze Zahlen, aber nicht sämtlich Null sind, ins Innere von $\mathfrak{R}$ fallen; denn sonst hätte $\mathfrak{R}$ mit dem Körper $\mathfrak{R}_{a,b,c}$ einen inneren Punkt gemein. Es sei nun $\mathfrak{S}$ eine Seitenfläche von $\mathfrak{R}$, die von $\mathfrak{S}$ und auch von der zu $\mathfrak{S}$ parallelen Seitenfläche verschieden ist, und $\frac{\bar{a}}{2}, \frac{\bar{b}}{2}, \frac{\bar{c}}{2}$ ein in dieser Seitenfläche, aber nicht auf ihrem Rande gelegener Punkt mit ganzen Zahlen $\bar{a}, \bar{b}, \bar{c}$; dann kann nicht $\bar{a} \equiv a, \bar{b} \equiv b, \bar{c} \equiv c \pmod{2}$ sein, da sonst $\frac{\bar{a}+a}{4}, \frac{\bar{b}+b}{4}, \frac{\bar{c}+c}{4}$ ein Punkt der eben besprochenen Art im Inneren von $\mathfrak{R}$ wäre. Da es nun im ganzen $2^3 - 1$ nach 2 inkongruente und von $0, 0, 0 \pmod{2}$ verschiedene Systeme $a, b, c \pmod{2}$ gibt, so besteht danach die Begrenzung von $\mathfrak{R}$ aus höchstens $2(2^3 - 1)$ Seitenflächen.

Die Lehrsätze I—VI sind hier nur für komplexe Polyeder im Raume von drei Dimensionen ausgesprochen, sie sind mit ihren hier auseinander-gesetzten Beweisen unmittelbar auf Mannigfaltigkeiten von beliebig vielen Veränderlichen zu übertragen.

Zürich, den 22. Juli 1897.

XXIII.

Über die Begriffe Länge, Oberfläche und Volumen.

(Referat über einen Vortrag: Jahresbericht der Deutschen Mathematiker-Vereinigung, Band IX, S. 115—121.)

1. Der Begriff des Ausdehnungsintegrals in einer Mannigfaltigkeit: $\int dx_1 dx_2 \dots dx_n$, d. i. für $n = 3$ der Begriff des *Volumens eines Körpers*, gehört zu den elementarsten Begriffen in der Analysis des Unendlichen; es knüpft dieser Begriff unmittelbar an den Begriff der Anzahl an. (Vgl. C. Jordan, Cours d'Analyse, 2^e éd., T. I, pp. 18—31.)

Wesentlich schwieriger als die Einführung des Volumenbegriffs ist die Begründung der *Länge einer Kurve* als Grenze der Länge von Polygonen, die der Kurve geeignet eingeschrieben sind, und der *Oberfläche einer krummen Fläche* als Grenze der Oberfläche von Polyedern, die in bezug auf die Fläche geeignet konstruiert sind.

Man kann jedoch diese anderen Begriffe Länge und Oberfläche auch allein aus dem Begriffe des Volumens mittels eines einfachen Grenzüberganges entwickeln:

Es sei C eine Kurve. Um jeden Punkt von C als Mittelpunkt denke man sich eine Kugel mit dem Radius r abgegrenzt, unter r eine feste positive Größe verstanden. Die Menge aller derjenigen Punkte des Raumes, welche in das Innere oder die Begrenzung von wenigstens einer dieser Kugel zu liegen kommen, definiert uns den *Bereich der Entfernung $\leq r$ von der Kurve C* . Es sei $V(r)$ das Volumen dieses Bereichs (falls ihm ein bestimmtes Volumen zukommt), so kann der Grenzwert von $\frac{V(r)}{\pi r^3}$ für ein nach Null abnehmendes r (falls dieser Grenzwert existiert), als die *Länge der Kurve C* eingeführt werden. — Es sei F eine Fläche. Man konstruiere in entsprechender Weise den *Bereich der Entfernung $\leq r$ von F* . Es sei $V(r)$ das Volumen dieses Bereichs, so kann der Grenzwert von $\frac{V(r)}{2r}$ für ein nach Null abnehmendes r (vorausgesetzt, daß die Größe $V(r)$ sowie dieser Grenzwert existiert), als die *Oberfläche der Fläche F* eingeführt werden.

Es ist einleuchtend, daß hierbei zunächst die Länge einer geradlinigen Strecke und die Oberfläche eines ebenen Dreiecks genau mit den gewöhnlich dafür angenommenen Werten sich ergeben; infolgedessen werden überhaupt in *regulären* Fällen Längen und Oberflächen in dem eben erklärten und andererseits in dem üblichen Sinne die gleichen Werte vorstellen.

2. Die soeben gegebene Definition einer Oberfläche führt uns zu einer bemerkenswerten Verallgemeinerung des Begriffs Oberfläche, indem wir an Stelle von Kugeln beliebige einander ähnliche und ähnlich gelegene konvexe Körper verwenden. Ich werde mich hier auf die Betrachtung geschlossener Flächen beschränken.

Unter einem *konvexen Körper* verstehe ich eine Punktmenge im Raume, welche abgeschlossen ist, die Eigenschaft hat, mit einer beliebigen Geraden stets entweder eine Strecke oder einen Punkt oder keinen Punkt gemein zu haben, und endlich nicht ganz in einer Ebene liegt. Eine *konvexe Fläche* bedeute die vollständige Begrenzung eines konvexen Körpers. Denken wir uns nun einen beliebigen konvexen Körper K zugrunde gelegt. Es sei G die Begrenzung von K und O irgendein bestimmter innerer Punkt von K . Ich will G die *Fläche der Distanz 1 von O* nennen (auch die *Eichfläche der Distanzen*). Ist P ein beliebiger Punkt und r eine positive GröÙe, so soll dann unter der *Fläche der Distanz r von P* diejenige konvexe Fläche H verstanden werden, welche P umschließt und mit G ähnlich und ähnlich gelegen ist derart, daß je zwei *gleichgerichtete* Radienvektoren von P nach H und von O nach G stets in ihren Längen das konstante Verhältnis $r:1$ darbieten. Der von dieser Fläche H umschlossene konvexe Körper ist dann der Bereich der Distanz $\leq r$ von P .

Es sei nun F eine beliebige, ganz im Endlichen gelegene *geschlossene Fläche*, d. h. eine Punktmenge, mittels deren der ganze Raum sich in zwei *abgeschlossene* Mengen, A und J , zerlegt, von denen eine jede die Menge F als *vollständige* Begrenzung besitzt und welche sonst untereinander keinen Punkt gemein haben; dasjenige von diesen zwei durch F geschiedenen Raumgebieten, in dem keine Grenzen für die Koordinaten der Punkte vorhanden sind, A , heiÙe der *äußere Raum* von F , das andere, J , der *innere Raum* von F . Wir denken uns um jeden Punkt von F den Körper der Distanz $\leq r$ von dem Punkte abgegrenzt. Es sei $Q_A(r)$, bzw. $Q_J(r)$ der Teil des Gebiets A , bzw. des Gebiets J , welcher von der Gesamtheit aller dieser Körper überdeckt wird, weiter $V_A(r)$, bzw. $V_J(r)$ das Volumen von $Q_A(r)$, bzw. von $Q_J(r)$, so heiÙe der Grenzwert von $\frac{V_A(r)}{r}$, bzw. $\frac{V_J(r)}{r}$ für ein nach Null abnehmendes r die *verallgemeinerte Außenoberfläche*, bzw. *Innenoberfläche* (abgekürzt v. A. O. bzw. v. J. O.) von F ,

immer stillschweigend vorausgesetzt, daß die betreffenden Volumina und Grenzen existieren. Die halbe Summe aus v. A.O. und v. J.O. von F heiße die *verallgemeinerte Oberfläche* von F .

3. Die Existenz der hier in Frage kommenden Grenzwerte läßt sich in einfacher Weise dartun, wenn die zu behandelnde Fläche F , ebenso wie die Eichfläche der Distanzen G , eine *konvexe Fläche* ist. Dann ist der innere Raum J von F ein konvexer Körper, und es sei C_0 sein Volumen. Der Raum $Q_A(r)$ wird hier jedesmal außer von F noch von einer zweiten konvexen Fläche $F_A(r)$ begrenzt, der Fläche derjenigen Punkte in A , für welche die kleinste Distanz von den Punkten in F gleich r ist.

Sind zunächst sowohl J wie K *Polyeder*, d. h. je von einer *endlichen* Anzahl von Ebenen vollständig begrenzt, so wird auch der von $F_A(r)$ begrenzte konvexe Körper bei beliebigem Werte des r ein Polyeder sein. Bei Zugrundelegung irgendeines Parallelkoordinatensystems erweisen sich alsdann die Koordinaten der Ecken von $F_A(r)$ als ganze lineare Funktionen von r , und vermöge der dreireihigen Determinanten für Volumina von Tetraedern erscheint hernach $V_A(r)$ als eine bestimmte ganze Funktion dritten Grades von r , d. h. *der vierte Differentialquotient der Funktion $V_A(r)$ ist gleich Null*.

Nun kann man eine beliebige konvexe Fläche stets durch zwei Polyederflächen annähern, von denen die eine ganz im inneren Raume, die andere ganz im äußeren Raume der Fläche verläuft, und welche miteinander ähnlich und ähnlich gelegen sind, und zwar noch derart, daß dabei das lineare Dilatationsverhältnis zur Erzeugung der zweiten Polyederfläche aus der ersten beliebig nahe an 1 liegt. (Vgl. meine Geometrie der Zahlen, S. 33.) Wenden wir diesen Hilfssatz sowohl in bezug auf die eine Fläche F , wie auch in bezug auf die Eichfläche G an, so zeigt sich, daß jene Eigenschaft des Verschwindens des vierten Differenzenquotienten von $V_A(r)$ sich von Polyederflächen sofort auf zwei beliebige konvexe Flächen F, G überträgt. Danach wird in allen Fällen das Volumen des von $F_A(r)$ begrenzten Körpers einen Ausdruck haben:

$$W(r) = C_0 + V_A(r) = C_0 + 3C_1r + 3C_2r^2 + C_3r^3,$$

wo C_0, C_1, C_2, C_3 gewisse von r unabhängige Konstanten sind.

Nunmehr wird $3C_1$ die verallgemeinerte Außenoberfläche von F .

Man erkennt leicht, daß bei Veränderung des Punktes O im Inneren von K die Größen C_0, C_1, C_2, C_3 sich nicht ändern, daß sie also nur von den zwei Flächen F und G , nicht von dem Punkte O abhängen. Vertauscht man die Rollen dieser zwei Flächen, so treten an die Stelle von $C_0, 3C_1, 3C_2, C_3$ die Werte $C_3, 3C_2, 3C_1, C_0$. Es ist also C_3 das Volumen von K und $3C_2$ die v. Außenoberfläche von G , wenn F als Eich-

fläche der Distanzen benutzt wird. Alle Größen C_0, C_1, C_2, C_3 sind danach positiv.

Für den hier eingeführten Begriff der v. Außenoberfläche heben wir als *in gewissem Sinne charakteristisch* die Eigenschaft hervor:

Enthält ein konvexer Körper einen anderen konvexen Körper in sich, so besitzt stets die Begrenzung des ersteren Körpers eine größere v. Außenoberfläche.

Weiter läßt sich zeigen: Die v. Innenoberfläche einer konvexen Fläche F ist gleich dem Werte, der für ihre v. Außenoberfläche entsteht, wenn die Eichfläche der Distanzen G durch die zu ihr in bezug auf den Punkt O symmetrische Fläche ersetzt wird. Danach erweisen sich v. Außenoberfläche und v. Innenoberfläche für eine konvexe Fläche F stets als gleich, wenn die Eichfläche eine *Fläche mit Mittelpunkt* ist.

Wird nunmehr als *Eichfläche eine Kugelfläche vom Radius 1* genommen, so erweist sich $3C_1$ als die *Oberfläche* der konvexen Fläche F im üblichen Sinne, während alsdann $3C_2$ die gesamte mittlere Krümmung von F darstellt.

Endlich machen wir die folgende Bemerkung:

Sind F und G miteinander ähnlich und ähnlich gelegen (worunter der Fall einzubegreifen ist, daß die Flächen durch bloße Parallelverschiebung auseinander hervorgehen), so ergibt sich

$$\frac{C_0}{C_1} = \frac{C_1}{C_2} = \frac{C_2}{C_3} = \sqrt[3]{\frac{C_0}{C_3}}.$$

4. Von diesen Betrachtungen will ich hier hauptsächlich Gebrauch machen, um einen *neuen und strengen Beweis* für den Satz zu geben, daß *unter allen konvexen Körpern von gleichem Volumen die Kugel die kleinste Oberfläche hat*, und um zugleich diesen Satz auf einen *inhaltsreicheren und analytisch einfacheren zurückzuführen*.

Ich stütze mich dabei auf den folgenden, von Herrn H. Brunn*) bewiesenen Satz:

Es seien J_0 und J_1 zwei beliebige konvexe Körper, die nicht miteinander sowohl ähnlich wie ähnlich gelegen sind, vom Volumen W_0 bzw. W_1 . Verbindet man jeden Punkt von J_0 mit jedem von J_1 und teilt die Verbindungsstrecke jedesmal in einem festen Verhältnisse $t:1-t$, wobei $0 < t < 1$ ist, so erfüllt die Menge aller verschiedenen solchen Teilpunkte wieder einen konvexen Körper J_t und gilt für dessen Volumen W_t die Ungleichung:

$$\sqrt[3]{W_t} > (1-t)\sqrt[3]{W_0} + t\sqrt[3]{W_1}.$$

*) Inauguraldissertation, München, 1887, S. 31. — Herr Brunn hat freilich an der angeführten Stelle ausdrücklich die Meinung geäußert: „Zum Beweise der Maximaleigenschaft der Kugel läßt sich dieser Satz nicht verwenden.“

Wir nehmen nun an, es seien die Flächen F und G nicht einander ähnlich und ähnlich gelegen, und können alsdann diesen Satz in der Weise anwenden, daß wir für J_0 und J_1 die zwei konvexen Körper nehmen, welche von zwei der oben betrachteten Flächen $F_A(r)$ für irgend zwei Werte $r = r_0$ und $r = r_1$ begrenzt werden; für J_t erscheint hierbei der von $F_A(r)$ für $r = (1-t)r_0 + tr_1$ begrenzte Körper. Die entstehende Ungleichung kommt nun darauf hinaus, daß die durch

$$w = \sqrt[3]{W(r)}$$

für $r \geq 0$ dargestellte Kurve, wenn man r und w als Abszisse und Ordinate in einer Ebene deutet, überall konvex auf ihrer der r -Achse abgewandten Seite ist, oder anders formuliert, daß

$$\frac{d^2 \sqrt[3]{W(r)}}{dr^2} < 0$$

ist im ganzen Bereiche $r \geq 0$.

Führen wir den in 3. gewonnenen Ausdruck von $W(r)$ ein, so muß danach

$$-\frac{1}{2} [W(r)]^{\frac{5}{3}} \frac{d^2 [W(r)]^{\frac{1}{3}}}{dr^2} = (C_1^2 - C_0 C_2) + (C_1 C_2 - C_0 C_3)r + (C_2^2 - C_1 C_3)r^2$$

für alle Werte $r \geq 0$ stets > 0 sein. Hierfür wieder sind die zwei Bedingungen

$$(I) \quad C_1^2 - C_0 C_2 > 0,$$

$$(II) \quad C_2^2 - C_1 C_3 > 0$$

oder also die Ungleichungen

$$\frac{C_0}{C_1} < \frac{C_1}{C_2} < \frac{C_2}{C_3}$$

erforderlich und hinreichend.

Beachten wir, daß wir die Rollen der beiden Flächen F und G vertauschen können und daß alsdann an Stelle der Größen C_0, C_1, C_2, C_3 diese Größen in umgekehrter Folge treten, so sehen wir, daß vermöge dieser Reziprozität die Ungleichung (I), für zwei beliebige konvexe Flächen genommen, bereits die Ungleichung (II) in sich schließt.

Nun leiten wir aus (I):

$$C_1^4 > C_0^2 C_2^2,$$

dann aus (II):

$$C_0^2 C_2^2 > C_0^2 C_1 C_3,$$

also mit Elimination von C_2 :

$$C_1^3 > C_0^2 C_3$$

her. Nehmen wir jetzt für die Eichfläche der Distanzen eine Kugelfläche vom Radius 1, so ist

$$C_3 = \frac{4\pi}{3},$$

ferner C_0 das Volumen und $3C_1$ die Oberfläche des von F begrenzten Körpers im gewöhnlichen Sinne; setzen wir

$$C_0 = \frac{4\pi}{3} R^3,$$

so folgt daher $3C_1 > 4\pi R^2$, d. i. der Satz, daß unter allen konvexen Körpern von gleichem Volumen die Kugel die kleinste Oberfläche besitzt.

Diese Eigenschaft der Kugel erscheint aber hier als Ausfluß des weit allgemeineren und analytisch einfacheren Theorems

$$C_1^2 > C_0 C_2,$$

welches sich auf zwei beliebige konvexe Körper bezieht. Dieses Theorem liefert im speziellen, wenn man für einen der Körper die Kugel nimmt, zwei neue, die Kugel unter allen konvexen Körpern charakterisierende Beziehungen: Nämlich unter allen konvexen Körpern von gleicher Oberfläche besitzt die Kugel erstens *die kleinste mittlere Krümmung*, zweitens *das größte Produkt aus Volumen und mittlerer Krümmung*. Aus beiden Sätzen zugleich resultiert als Folgerung jene bekannte isoperimetrische Eigenschaft der Kugel.

5. Der Schluß des Vortrags brachte noch ein Theorem über die Bestimmung einer geschlossenen konvexen Fläche, wenn für sie in jedem Punkte die Gaußsche Krümmung als Funktion der Normalenrichtung in dem Punkte beliebig vorgeschrieben ist.

Über die geschlossenen konvexen Flächen.*)

Im folgenden teile ich einige Resultate einer Untersuchung über geschlossene konvexe Flächen mit. Ich bin zu dieser Untersuchung durch den Satz, daß unter allen Körpern gleichen Volumens die Kugel die kleinste Oberfläche hat, angeregt worden.

Der Einfachheit wegen will ich mich hier auf die Betrachtung solcher konvexer Flächen beschränken, die in jedem Punkte eine bestimmte Tangentialebene und bestimmte endliche Hauptkrümmungsradien haben.

1. Es bedeute Ω die Kugelfläche mit dem Nullpunkte O als Mittelpunkt vom Radius 1, und es seien $\cos \psi \sin \vartheta$, $\sin \psi \sin \vartheta$, $\cos \vartheta$ die Koordinaten eines beliebigen Punktes Π dieser Kugelfläche.

Es bedeute K einen beliebigen konvexen Körper; es sei

$$x \cos \psi \sin \vartheta + y \sin \psi \sin \vartheta + z \cos \vartheta = H$$

die Gleichung derjenigen Tangentialebene an die Oberfläche von K , welche die Richtung $O\Pi$ als äußere Normale hat. Dabei stellt dann H eine eindeutige Funktion auf der Kugelfläche Ω vor, und durch Angabe dieser Funktion $H(\vartheta, \psi)$ ist der Körper K bereits vollständig bestimmt.

Setzt man

$$R = \frac{\partial^2 H}{\partial \vartheta^2} + H, \quad S = \frac{1}{\sin \vartheta} \frac{\partial^2 H}{\partial \vartheta \partial \psi} - \frac{\cos \vartheta}{\sin^2 \vartheta} \frac{\partial H}{\partial \psi},$$

$$T = \frac{1}{\sin^2 \vartheta} \frac{\partial^2 H}{\partial \psi^2} + \frac{\cos \vartheta}{\sin \vartheta} \frac{\partial H}{\partial \vartheta} + H,$$

so ist $Ru^2 + 2Suv + Tv^2$ eine positive quadratische Form. Bekanntlich ist $R + T$ die Summe, $RT - S^2$ das Produkt der Hauptkrümmungsradien im Berührungspunkte jener Tangentialebene.

*) Diese Abhandlung erschien in den Comptes rendus de l'Académie des Sciences, Paris 1901, t. 132, pp. 21—24, in einer vom Verfasser herrührenden französischen Übersetzung unter dem Titel: Sur les surfaces convexes fermés. Dieser Abdruck folgt dem deutschen Originalmanuskript. Einige Zusätze der französischen Ausgabe sind übersetzt und in Klammern [] hinzugefügt worden. (Anm. d. Herausg.)

2. Es seien nun K_1, K_2, K_3 irgend drei konvexe Körper, und die Größen H, R, S, T für sie mögen durch Hinzufügen von Indizes [bzw. 1, 2, 3] unterschieden werden.

Ich bezeichne alsdann den Wert des Integrals

$$\frac{1}{6} \int H_1 (R_2 T_3 - 2 S_2 S_3 + T_2 R_3) d\omega = A_{123},$$

erstreckt über alle Elemente $d\omega = \sin \vartheta d\vartheta d\psi$ der Kugelfläche Ω , als das *gemischte Volumen* der Körper K_1, K_2, K_3 .

Dieses Integral fällt stets positiv aus, und *der Wert dieses Integrals ändert sich nicht, wenn die Körper irgendwie permutiert werden*, ferner auch nicht, wenn die Körper irgendwelchen Translationen unterworfen werden.

Sind die drei Körper identisch mit einem einzigen Körper, so stellt das Integral das Volumen dieses Körpers vor.

3. Drei konvexe Körper K_1, K_2, K_3 mögen *unabhängig* heißen, wenn zwischen ihren Funktionen H_1, H_2, H_3 keine identische Relation

$$w_1 H_1 + w_2 H_2 + w_3 H_3 = x_0 \cos \psi \sin \vartheta + y_0 \sin \psi \sin \vartheta + z_0 \cos \vartheta$$

mit irgendwelchen 6 [von ϑ und ψ unabhängigen] Konstanten $w_1, w_2, w_3, x_0, y_0, z_0$ besteht.

Sind K_1, K_2, K_3 beliebige konvexe Körper mit den drei zugehörigen Funktionen H_1, H_2, H_3 und sind w_1, w_2, w_3 irgendwelche Konstanten ≥ 0 , jedoch nicht alle drei gleich Null, so wird durch die Funktion

$$H = w_1 H_1 + w_2 H_2 + w_3 H_3$$

jedesmal wieder ein konvexer Körper K bestimmt. Das Volumen dieses Körpers ist alsdann

$$F(w_1, w_2, w_3) = \sum_i \sum_k \sum_l A_{ikl} w_i w_k w_l \quad (i, k, l = 1, 2, 3),$$

wo A_{ikl} das gemischte Volumen von K_i, K_k, K_l ist. Nunmehr gilt folgendes Theorem:

Die Fläche

$$F(w_1, w_2, w_3) = 1,$$

[wo w_1, w_2, w_3 als rechtwinklige Koordinaten aufgefaßt werden] ist im Gebiete $w_1 \geq 0, w_2 \geq 0, w_3 \geq 0$ eine konvexe Fläche, welche ihre Konvexität nach der dem Nullpunkte abgewandten Seite kehrt.

Sind K_1, K_2, K_3 *unabhängig*, so enthält diese Fläche auch niemals eine geradlinige Strecke. Als dann ist die Determinante

$$\begin{vmatrix} A_{111} & A_{112} & A_{113} \\ A_{121} & A_{122} & A_{123} \\ A_{131} & A_{132} & A_{133} \end{vmatrix} > 0$$

und sind die zwei entsprechenden Determinanten [die als ersten Index 2 oder 3 an Stelle von 1 enthalten] positiv und sind

$$\begin{vmatrix} A_{111}, A_{112} \\ A_{121}, A_{122} \end{vmatrix} < 0$$

und alle entsprechenden Verbindungen sämtlich negativ.

Insbesondere folgt daraus:

Sind zwei konvexe Körper K_1, K_2 nicht ähnlich und ähnlich gelegen, so gilt stets

$$A_{111}A_{122} < A_{112}^2, A_{112}A_{222} < A_{122}^2.$$

Hierin ist A_{111} das Volumen von K_1 . Nimmt man nun insbesondere K_2 gleich einer Kugel vom Radius 1, so wird ferner $3A_{112}$ die Größe der Oberfläche von K_1 und weiter $3A_{122}$ die gesamte mittlere Krümmung dieser Fläche, während endlich $A_{222} = \frac{4\pi}{3}$ ist. Auf diese Weise erhält man die Sätze:

Unter allen konvexen Körpern von gleicher Oberfläche hat die Kugel das größte Produkt aus Volumen und mittlerer Krümmung, ferner die kleinste mittlere Krümmung und durch beides schließlich das größte Volumen.

Eine andere Folgerung der letzten Ungleichungen ist diese:

Hat ein konvexer Körper ein Volumen = 1, ohne ein Würfel mit Seitenflächen parallel den Koordinatenebenen zu sein, so ist stets das arithmetische Mittel aus den Flächeninhalten seiner drei Projektionen auf die drei Koordinatenebenen > 1 .

4. Es sei $G(\vartheta, \psi)$ eine auf der Kugelfläche Ω eindeutige und stetige Funktion, welche daselbst überall > 0 ist, und noch derart beschaffen ist, daß die drei Integrale

$$\int \frac{\cos \psi \sin \vartheta}{G} d\omega, \quad \int \frac{\sin \psi \sin \vartheta}{G} d\omega, \quad \int \frac{\cos \vartheta}{G} d\omega,$$

über diese ganze Kugelfläche erstreckt, sämtlich verschwinden. Alsdann existiert stets ein konvexer Körper K , für den die Gaußsche Krümmung in dem Punkte, dessen äußere Normale die Richtungskosinus $\cos \psi \sin \vartheta, \sin \psi \sin \vartheta, \cos \vartheta$ hat, gleich $G(\vartheta, \psi)$ ist, und dieser Körper ist bis auf eine beliebige Translation, die man ihm noch erteilen kann, eindeutig bestimmt.

Man kann nämlich unter allen konvexen Körpern vom Volumen 1 zunächst einen Körper derart bestimmen, daß für seine Funktion H der Wert des Integrals

$$J = \int \frac{H}{G} d\omega$$

möglichst klein ausfällt. Dieser Körper ist bis auf eine Translation völlig bestimmt; erlangt für ihn das Integral J den Wert λ , so entsteht aus diesem Körper durch Dilatation im Verhältnis $\sqrt{\lambda}:1$ der gesuchte Körper K , für den $RT - S^2 = \frac{1}{G}$ ist.

Theorie der konvexen Körper, insbesondere Begründung ihres Oberflächenbegriffs. *)

I. Kapitel.

Die konvexen Körper als Punktgebilde.

§ 1. Definition der konvexen Körper.

Eine Punktmenge soll ein *konvexer Körper* heißen, wenn sie 1^o mit einer beliebigen Geraden jedesmal sei es eine endliche Strecke, sei es einen Punkt, sei es keinen Punkt, gemein hat und 2^o nicht ganz in einer Ebene gelegen ist.

Es sei $\mathcal{K}$ ein konvexer Körper. Nach Voraussetzung enthält $\mathcal{K}$ irgend vier nicht in einer Ebene gelegene Punkte a, b, c, d . Mit a und b enthält $\mathcal{K}$ nach Voraussetzung von der durch a und b gehenden Geraden jedenfalls die ganze Strecke ab , sodann mit c jeden Punkt des Dreiecks abc , weiter mit d jeden Punkt des Tetraeders $abcd$. Es besitzt also die Punktmenge $\mathcal{K}$ jedenfalls auch *innere* Punkte. — Indem wir eine Parallelverschiebung (Translation) von $\mathcal{K}$ zulassen, können wir einen beliebig gegebenen Punkt als inneren von $\mathcal{K}$ annehmen.

§ 2. Die Distanzfunktion für einen konvexen Körper, welcher den Nullpunkt im Inneren enthält.

Der konvexe Körper $\mathcal{K}$ enthalte den Anfangspunkt o der rechtwinkligen Koordinaten x, y, z als einen *inneren* Punkt.

Ziehen wir vom Nullpunkte o aus in einer beliebigen Richtung einen geradlinigen Strahl, so muß dieser nach 1^o mit $\mathcal{K}$ jedesmal eine bestimmte Strecke op_0 gemein haben. Es seien x_0, y_0, z_0 die Koordinaten des Endpunktes p_0 dieser Strecke, so wollen wir, wenn x, y, z die Koordinaten

*) Diese bisher unveröffentlichte Abhandlung, die sich im Nachlaß gefunden hat, ist der erste Teil eines größeren Werkes über die Theorie der konvexen Körper. Vom zweiten Teil sind nur wenige Paragraphen ausgeführt, deren Resultate, wenn auch nach andern Methoden abgeleitet, in die Abhandlung „Volumen und Oberfläche“ übergegangen sind. (Anm. d. Herausg.)

eines völlig beliebigen Punktes p jenes Strahles sind, den gemeinsamen Wert der Quotienten

$$\frac{x}{x_0} = \frac{y}{y_0} = \frac{z}{z_0} = f(x, y, z)$$

setzen und die so entstehende Funktion $f(x, y, z)$ die *Distanzfunktion* des Körpers $\mathfrak{K}$ nennen.

Für die verschiedenen Punkte p_0 ist dann $f(x_0, y_0, z_0) = 1$, und für die Punkte der Menge $\mathfrak{K}$ und nur für diese Punkte gilt $f(x, y, z) \leq 1$. Die hier eingeführte Funktion $f(x, y, z)$ besitzt nun die folgenden Eigenschaften:

1. Für jedes von $0, 0, 0$ verschiedene Wertsystem x, y, z ist $f(x, y, z) > 0$; ferner ist $f(0, 0, 0) = 0$;

2. Man hat immer

$$f(tx, ty, tz) = tf(x, y, z),$$

wenn $t > 0$ ist.

3. Für beliebige zwei Systeme $x_1, y_1, z_1; x_2, y_2, z_2$ gilt allgemein:

$$f(x_1, y_1, z_1) + f(x_2, y_2, z_2) \geq f(x_1 + x_2, y_1 + y_2, z_1 + z_2).$$

Ist eines der Systeme mit $0, 0, 0$ identisch, so versteht sich diese Relation wegen $f(0, 0, 0) = 0$ von selbst. Wenn $\frac{x_2}{x_1} = \frac{y_2}{y_1} = \frac{z_2}{z_1} > 0$ ist, folgt sie aus 2. Nehmen wir jetzt an, es seien die Punkte $p_1(x_1, y_1, z_1)$ und $p_2(x_2, y_2, z_2)$ von o verschieden und auch die Richtungen op_1 und op_2 voneinander verschieden. Wir setzen

$$f(x_1, y_1, z_1) + f(x_2, y_2, z_2) = s, \quad f(x_1, y_1, z_1) = ts, \quad f(x_2, y_2, z_2) = (1 - t)s;$$

dabei ist $0 < t < 1$. Es sei q_1 der Punkt mit den Koordinaten $\frac{x_1}{ts}, \frac{y_1}{ts}, \frac{z_1}{ts}$ und q_2 der Punkt mit den Koordinaten $\frac{x_2}{(1-t)s}, \frac{y_2}{(1-t)s}, \frac{z_2}{(1-t)s}$; für den einen wie für den anderen Punkt wird dann nach 2.: $f(x, y, z) = 1$. Es gehören diese Punkte also zur Menge $\mathfrak{K}$ und muß infolgedessen auch jeder Punkt der Strecke q_1q_2 zu $\mathfrak{K}$ gehören, insbesondere also der Punkt $tq_1 + (1-t)q_2$ (d.h. der Schwerpunkt von q_1 und q_2 , wenn dem Punkte q_1 die Masse t , dem Punkte q_2 die Masse $1 - t$ beigelegt wird). Dieser Punkt hat die Koordinaten $\frac{x_1 + x_2}{s}, \frac{y_1 + y_2}{s}, \frac{z_1 + z_2}{s}$ und folgt nunmehr für ihn $f(x, y, z) \leq 1$, d. i. mit Rücksicht auf 2. die Ungleichung in 3.

Die Beziehung $f(-x, -y, -z) = f(x, y, z)$ wird dann und nur dann statthaben, wenn der Körper $\mathfrak{K}$ den Punkt o als Mittelpunkt hat.

§ 3. Stetigkeit der Distanzfunktion und Folgerungen.

Aus 3. folgt

$$f(x, y, z) \leq f(x, y, 0) + f(0, 0, z) \leq f(x, 0, 0) + f(0, y, 0) + f(0, 0, z).$$

Es werde der größte Wert unter den sechs Größen $f(\pm 1, 0, 0), f(0, \pm 1, 0),$

$f(0, 0, \pm 1)$ mit G bezeichnet, so folgt hieraus mit Rücksicht auf die Regel 2. in § 2 weiter:

$$(1) \quad f(x, y, z) \leq G(|x| + |y| + |z|).$$

Das durch $|x| + |y| + |z| \leq \frac{1}{G}$ definierte *Oktäeder* und um so mehr der durch

$$|x| \leq \frac{1}{3G}, \quad |y| \leq \frac{1}{3G}, \quad |z| \leq \frac{1}{3G}$$

definierte *Würfel* sind danach ganz im Körper $\mathfrak{K}$ enthalten.

Nun geht aus der Regel 3. in § 2:

$$(2) \quad -f(-x_2, -y_2, -z_2) \leq f(x_2 + x_1, y_2 + y_1, z_2 + z_1) - f(x_1, y_1, z_1) \leq f(x_2, y_2, z_2),$$

also

$$(3) \quad |f(x_2 + x_1, y_2 + y_1, z_2 + z_1) - f(x_1, y_1, z_1)| \leq G(|x_2| + |y_2| + |z_2|) \\ \leq \sqrt{3} G \sqrt{x_2^2 + y_2^2 + z_2^2}$$

hervor. Diese Beziehung zeigt, daß $f(x, y, z)$ eine *stetige* Funktion der Argumente x, y, z ist.

Wir wollen stets mit $\mathfrak{E}$ die Kugelfläche vom Radius 1 mit dem Nullpunkt $\mathfrak{o}$ als Mittelpunkt bezeichnen, also den Bereich derjenigen Punkte α, β, γ , für welche $\alpha^2 + \beta^2 + \gamma^2 = 1$ ist. Unter einer *Richtung* (α, β, γ) , wobei dann stets α, β, γ die Koordinaten eines Punktes $\mathfrak{r}$ dieser Kugelfläche $\mathfrak{E}$ bedeuten sollen, werden wir die Richtung von $\mathfrak{o}$ nach $\mathfrak{r}$ verstehen.

Der Ausdruck $f(\alpha, \beta, \gamma)$ wird als stetige Funktion seiner Argumente α, β, γ in dem *abgeschlossenen* Bereiche der Kugelfläche $\mathfrak{E}$ ein bestimmtes *Minimum* besitzen. Der Wert dieses Minimum wird wie alle Werte, welche $f(x, y, z)$ außerhalb des Nullpunktes hat, wesentlich positiv sein. Wir bezeichnen ihn mit $\frac{1}{R}$, wobei dann R eine bestimmte positive endliche Größe sein wird. Ferner wird $f(\alpha, \beta, \gamma)$ auf $\mathfrak{E}$ ein bestimmtes *Maximum* haben, das wir $= \frac{1}{r}$ setzen. Dabei gilt jedenfalls $G \leq \frac{1}{r}$. Wir haben nun auf $\mathfrak{E}$ stets $\frac{1}{r} \geq f(\alpha, \beta, \gamma) \geq \frac{1}{R}$. Auf einer Kugelfläche von einem Radius t mit dem Mittelpunkt $\mathfrak{o}$ wird dann stets $\frac{t}{r} \geq f(x, y, z) \geq \frac{t}{R}$ sein, d. h. es gilt allgemein

$$(4) \quad \frac{1}{r} \sqrt{x^2 + y^2 + z^2} \geq f(x, y, z) \geq \frac{1}{R} \sqrt{x^2 + y^2 + z^2}.$$

Ziehen wir zunächst die untere Grenze für $f(x, y, z)$ hier in Betracht, so enthält danach die Kugel vom Radius R und um so mehr der durch

$$|x| \leq R, \quad |y| \leq R, \quad |z| \leq R$$

definierte Würfel den Körper $\mathfrak{K}$ ganz in sich. Insbesondere ist hiernach ein konvexer Körper stets *ganz im Endlichen gelegen*, das soll heißen, es bestehen stets für die Koordinaten seiner Punkte obere und untere Grenzen.

Wegen der Stetigkeit der Funktion $f(x, y, z)$ ist der konvexe Körper $\mathfrak{K}$ stets eine *abgeschlossene* Punktmenge, d. h. alle *Häufungsstellen* dieser Menge sind in ihr selbst enthalten, und wird die *Begrenzung* dieser Menge genau von denjenigen Punkten gebildet, für welche $f(x, y, z) = 1$ ist. Die vollständige Begrenzung eines konvexen Körpers bezeichnen wir als eine *konvexe Fläche*.

Derjenige Punkt x, y, z der Begrenzung von $\mathfrak{K}$, welcher in einer bestimmten Richtung (α, β, γ) von o aus liegt, wird bestimmt durch

$$\frac{x}{\alpha} = \frac{y}{\beta} = \frac{z}{\gamma} = \frac{f(x, y, z)}{f(\alpha, \beta, \gamma)} \text{ und } f(x, y, z) = 1,$$

also ist $\frac{1}{f(\alpha, \beta, \gamma)}$ seine Entfernung vom Nullpunkte.

Die Eigenschaft 1^o (§ 1) einer Punktmenge $\mathfrak{K}$, mit jeder Geraden eine Strecke oder einen Punkt oder keinen Punkt gemein zu haben, können wir in vielen Fällen vorteilhafter durch die Gesamtheit der folgenden drei Eigenschaften ersetzen:

1a) Mit irgend zwei Punkten der Menge gehört stets auch jeder Punkt der dieselben verbindenden Strecke zur Menge,

1b) die Menge ist ganz im Endlichen gelegen,

1c) sie ist eine abgeschlossene Punktmenge.

Daß die Eigenschaft 1^o bei einer Punktmenge, die nicht ganz in einer Ebene liegt, diese Eigenschaften 1a), 1b), 1c) nach sich zieht, haben wir soeben gesehen. Indem wir zu diesem Satze für den Raum die analogen Tatsachen für die Ebene und die gerade Linie hinzunehmen, erkennen wir ganz allgemein, daß aus der Eigenschaft 1^o die Eigenschaften 1a), 1b), 1c) folgen.

Setzen wir nun umgekehrt für eine Punktmenge $\mathfrak{K}$ die drei letzteren Eigenschaften voraus. Betrachten wir irgendeine Gerade, die wenigstens zwei Punkte von $\mathfrak{K}$ enthält, und bestimmen wir die verschiedenen Punkte auf der Geraden mittels einer Kartesischen Koordinate. Wegen 1b) haben wir dann die Werte dieser Koordinate bei allen denjenigen Punkten der Geraden, welche zu $\mathfrak{K}$ gehören, eine obere und eine untere Grenze. Wegen 1c) gehören die diesen Grenzwerten der Koordinate entsprechenden zwei Punkte der Geraden dann ebenfalls selbst zu $\mathfrak{K}$, und wegen 1a) hat dann $\mathfrak{K}$ mit der Geraden genau die diese zwei Grenzpunkte verbindende Strecke gemein.

Nunmehr können wir die bisher erlangten Resultate in folgender Weise umkehren:

Ist $f(x, y, z)$ eine beliebige Funktion mit den in § 2 bei 1., 2., 3. angegebenen Eigenschaften, so stellt der durch $f(x, y, z) \leq 1$ definierte Bereich $\mathfrak{K}$ stets einen konvexen Körper mit dem Nullpunkte als inneren Punkt vor.

Denn die betreffende Funktion $f(x, y, z)$ ist jedesmal stetig und danach der Nullpunkt, für den $f(0, 0, 0) = 0$ ist, ein innerer Punkt von $\mathfrak{R}$, die Menge $\mathfrak{R}$ also gewiß nicht ganz in einer Ebene gelegen, ferner ist dann $\mathfrak{R}$ ganz im Endlichen gelegen und abgeschlossen. Endlich gilt noch die Eigenschaft 1a). Denn sind $p_1(x_1, y_1, z_1)$ und $p_2(x_2, y_2, z_2)$ irgend zwei verschiedene Punkte aus $\mathfrak{R}$, also dafür $f(x_1, y_1, z_1) \leq 1$, $f(x_2, y_2, z_2)$ und ist t ein Wert > 0 und < 1 , so folgt für den Punkt $tp_1 + (1-t)p_2$ der Strecke p_1p_2 aus § 2, 2. und 3.:

$f(tx_1 + (1-t)x_2, ty_1 + (1-t)y_2, tz_1 + (1-t)z_2) \leq t + (1-t) = 1$,
es gehört also die ganze, p_1 und p_2 verbindende Strecke zu $\mathfrak{R}$.

§ 4. Polyeder. Stützebenen.

Bedeutet p einen Punkt mit den Koordinaten x, y, z und s eine Konstante, so soll der Punkt, dessen Koordinaten die Werte sx, sy, sz haben, mit sp bezeichnet werden. Sind $p'(x', y', z')$ und $p''(x'', y'', z'')$ zwei Punkte, so soll unter $p' + p''$ der Punkt mit den Koordinaten $x' + x'', y' + y'', z' + z''$ verstanden werden. — Bedeutet $\mathfrak{R}$ eine Menge von Punkten p und s eine Konstante, so soll $s\mathfrak{R}$ die Menge der Punkte sp vorstellen.

Sind $p_1, p_2, \dots, p_n$ eine *endliche* Anzahl von Punkten, so bildet die Gesamtheit aller derjenigen Punkte p , welche sich in der Form

$$(5) \quad p = t_1 p_1 + t_2 p_2 + \dots + t_n p_n$$

darstellen, so daß dabei

$$(6) \quad t_1 \geq 0, t_2 \geq 0, \dots, t_n \geq 0, t_1 + t_2 + \dots + t_n = 1$$

ist, einen Bereich, der die Eigenschaften 1a), 1b), 1c) eines konvexen Körpers besitzt, und den wir als den *konvexen Bezirk* ($p_1, p_2, \dots, p_n$) bezeichnen wollen. Der Bereich enthält die Punkte $p_1, p_2, \dots, p_n$ und muß in jedem konvexen Körper, der diese Punkte sämtlich aufnimmt, stets vollständig enthalten sein.

Besitzt einer jener Punkte, z. B. p_n außer der selbstverständlichen Darstellung $t_1 = 0, \dots, t_{n-1} = 0, t_n = 1$ noch eine zweite Darstellung in der Form (5) und (6), wobei dann also $t_n < 1$ ist, so erweist sich dadurch p_n bereits als Punkt des konvexen Bezirks ($p_1, p_2, \dots, p_{n-1}$) und ist also der konvexe Bezirk ($p_1, p_2, \dots, p_{n-1}, p_n$) in dem konvexen Bezirke ($p_1, p_2, \dots, p_{n-1}$) enthalten und daher mit letzterem Bezirke identisch. Indem wir in solcher Weise, soweit als möglich, die Anzahl der dem konvexen Bezirk zugrunde liegenden Punkte verringern, verbleiben schließlich gewisse Punkte jener Reihe, etwa $p_1, p_2, \dots, p_m$ ($m \leq n$) derart, daß der konvexe Bezirk ($p_1, p_2, \dots, p_n$) noch mit ($p_1, p_2, \dots, p_m$) identisch ist, daß aber jeder dieser Punkte $p_1, p_2, \dots, p_m$ in der Form

$$t_1 p_1 + t_2 p_2 + \dots + t_m p_m \text{ mit } t_1 \geq 0, t_2 \geq 0, \dots, t_m \geq 0, t_1 + t_2 + \dots + t_m = 1$$

und daher auch in der ursprünglich betrachteten Form (5) und (6) nur auf eine Weise darstellbar ist.

Diejenigen Punkte p_h der Reihe $p_1, p_2, \dots, p_n$, welche in der Form (5) und (6) nur so darzustellen sind, daß $t_h = 1$ und die anderen Größen t_g ($g \neq h$) Null sind, heißen die *Eckpunkte* des konvexen Bezirks $(p_1, p_2, \dots, p_n)$. Auf einer Geraden durch einen Eckpunkt p_h können niemals zu beiden Seiten von p_h Punkte des konvexen Bezirks liegen.

Liegen die Punkte $p_1, p_2, \dots, p_n$ nicht sämtlich in einer Ebene, so ist der konvexe Bezirk $(p_1, p_2, \dots, p_n)$ ein konvexer Körper und heißt ein (konvexes) *Polyeder*. — Wir bemerken, daß alsdann jeder Punkt p , der in der Form (5) und (6) mit lauter *positiven* Faktoren $t_1, t_2, \dots, t_n$ erscheint, stets ein *innerer* Punkt des Polyeders ist. Denn bedeutet e irgendeinen inneren Punkt des Polyeders, wobei nun nicht $e = p$ sei, so erscheinen die Punkte $(1+t)p - te$ für hinreichend kleines positives t , d. s. also Punkte auf dem Strahle von e durch p über p hinaus, ebenfalls noch in der Form (5) und (6), während diese Punkte nach § 1 außerhalb des Polyeders liegen müßten, wenn p ein Punkt seiner Begrenzung wäre.

Liegen $p_1, p_2, \dots, p_n$ sämtlich in einer Ebene, aber nicht sämtlich auf einer Geraden, so heißt der konvexe Bezirk $(p_1, p_2, \dots, p_n)$ ein (konvexes) *Polygon*. Ein jeder Punkt p der Form (5) und (6) mit lauter *positiven* Koeffizienten t_h ist dann stets ein *innerer* Punkt des Polygons in der Mannigfaltigkeit seiner Ebene, oder wie wir anstatt dessen lieber sagen wollen, da ein Polygon als eine Punktmenge im Raume überhaupt keinen inneren Punkt hat, ein *inwendiger* Punkt des Polygons.

Liegen $p_1, p_2, \dots, p_n$ sämtlich auf einer Geraden, so reduziert sich der konvexe Bereich $(p_1, p_2, \dots, p_n)$ auf eine *Strecke*, wofern er nicht bloß einen einzelnen *Punkt* vorstellt.

Ist

$$(7) \quad \alpha x + \beta y + \gamma z = d,$$

wobei $\alpha^2 + \beta^2 + \gamma^2 = 1$ statthabe, die Gleichung einer Ebene, so bezeichnen wir diejenigen Punkte x, y, z , für welche

$$\alpha x + \beta y + \gamma z > d$$

gilt, als auf der *Seite* (α, β, γ) dieser Ebene gelegen.

Es sei $\mathfrak{R}$ eine abgeschlossene Punktmenge; finden wir in (7) eine Ebene, welche wenigstens einen Punkt von $\mathfrak{R}$ enthält, aber auf der Seite (α, β, γ) von sich keinen Punkt von $\mathfrak{R}$ liegen hat, so daß also in $\mathfrak{R}$ durchweg

$$(8) \quad \alpha x + \beta y + \gamma z \leq d$$

gilt, so bezeichnen wir eine solche Ebene als *Stützebene* an $\mathfrak{R}$, mit der

(äußeren) *Normalen* (α, β, γ) und mit der *Bedingung* (8). Jeder Punkt von $\mathfrak{K}$ in der Stützebene gehört notwendig der Begrenzung von $\mathfrak{K}$ an.

Wir beweisen jetzt den folgenden Satz:

Für ein Polyeder kann man stets eine endliche Anzahl von Stützebenen angeben, derart, daß durch die Bedingungen dieser Stützebenen der Bereich des Polyeders vollständig charakterisiert ist.

Es sei $\mathfrak{P}$ ein Polyeder mit den Eckpunkten $p_1, p_2, \dots, p_m$. Es sei e irgendein innerer Punkt von $\mathfrak{P}$. Wir konstruieren jede Verbindungsstrecke $p_i p_j$ von zwei der Eckpunkte und jedesmal, wenn diese Strecke nicht e enthält, die Ebene durch e, p_i, p_j . Es sei nun p ein solcher Punkt der *Begrenzung* von $\mathfrak{P}$, der in keiner dieser Ebenen $ep_i p_j$ und daher gewiß auch auf keiner Strecke $p_i p_j$ liegt. Der Punkt p erscheint irgendwie in einer Form

$$p = t_i p_i + t_j p_j + \dots + t_l p_l,$$

so daß $i, j, \dots, l$ Werte der Reihe $1, 2, \dots, m$ und $t_i, t_j, \dots, t_l$ sämtlich > 0 und dabei $t_i + t_j + \dots + t_l = 1$ ist. Dabei können die Punkte nicht sämtlich auf einer Geraden liegen, da ja p auf keiner Strecke $p_i p_j$ liegt; ihre Anzahl ist also ≥ 3 . Andererseits aber müssen alle diese Punkte in einer Ebene liegen, denn sonst würde der konvexe Bezirk $(p_i, p_j, \dots, p_l)$ ein Polyeder und nach einer oben gemachten Bemerkung p ein innerer Punkt dieses Polyeders und daher um so mehr ein innerer Punkt von $\mathfrak{P}$ selbst sein. Also bildet der konvexe Bereich $(p_i, p_j, \dots, p_l)$ ein *Polygon*, und ist p nach einer zweiten Bemerkung oben ein *inwendiger* Punkt dieses Polygons.

Jetzt muß die Ebene dieses Polygons eine Stützebene an das Polyeder $\mathfrak{P}$ sein. Denn die Ebene enthält gewiß nicht e , von den Eckpunkten $p_1, p_2, \dots, p_m$ befindet sich dann notwendig wenigstens einer, etwa p_g , auf derselben Seite dieser Ebene wie e . Hätte nun die Ebene auch auf der entgegengesetzten Seite einen Eckpunkt, etwa p_h liegen, so würden die beiden konvexen Bezirke $(p_g, p_i, p_j, \dots, p_l)$ und $(p_h, p_i, p_j, \dots, p_l)$ (die *Pyramiden* mit jenem Polygon als Basis und p_g bzw. p_h als Spitze) durch diese Ebene getrennt sein, und p würde als inwendiger Punkt ihrer gemeinsamen Basis ein innerer Punkt in dem aus den beiden Pyramiden zusammengesetzten Bereiche und um so mehr ein innerer Punkt in $\mathfrak{P}$ sein, was gegen die Voraussetzung wäre. Damit sind wir zunächst zu folgendem Ergebnisse gekommen:

Bestimmen wir alle Ebenen durch je drei nicht in einer Geraden gelegene der Eckpunkte $p_1, p_2, \dots, p_m$, so haben wir in dieser endlichen Anzahl von Ebenen gewisse Stützebenen an $\mathfrak{P}$. Es seien

$$(9) \quad \alpha^{(x)} x + \beta^{(x)} y + \gamma^{(x)} z \leq d^{(x)} \quad (x = 1, 2, \dots, \nu)$$

die Bedingungen dieser verschiedenen Stützebenen an $\mathfrak{P}$, so gehen diese Stützebenen jedenfalls durch jeden solchen Punkt p der Begrenzung von $\mathfrak{P}$, der in keiner der Ebenen ep, p_j liegt. Nun können wir aber eine Richtung $(\alpha^*, \beta^*, \gamma^*)$, die in einer jener Ebenen ep, p_j auftritt, da die Anzahl der Ebenen endlich ist, jedenfalls als eine Häufungsstelle von solchen Richtungen (α, β, γ) darstellen, die nicht diesen Ebenen angehören, und danach ist weiter jeder solche Punkt p^* der Begrenzung von $\mathfrak{P}$, für den ep^* in eine der Ebenen ep, p_j fällt, eine Häufungsstelle von solchen Punkten p der Begrenzung, für die ep nicht in diese Ebenen fällt, welche also in den Ebenen (9) liegen. Die Ebenen (9) müssen deshalb überhaupt *jeden* Punkt der Begrenzung von $\mathfrak{P}$ aufnehmen.

Es gelten im ganzen Bereiche $\mathfrak{P}$ die Ungleichungen (9). Ist aber q ein Punkt außerhalb $\mathfrak{P}$, so enthält die Strecke eq einen Punkt p der Begrenzung von $\mathfrak{P}$. Führt durch diesen Punkt p etwa die κ^{te} der Stützebenen (9), so ist der Ausdruck

$$\alpha^{(\kappa)}x + \beta^{(\kappa)}y + \gamma^{(\kappa)}z$$

für e kleiner als $d^{(\kappa)}$, für p gleich $d^{(\kappa)}$, für q daher größer als $d^{(\kappa)}$, also erfüllt q nicht die Bedingungen (9). Danach ist der Bereich $\mathfrak{P}$ genau durch die Ungleichungen (9) charakterisiert.

In jeder der Ebenen (9) gehört zu $\mathfrak{P}$ ein gewisses Polygon, das wir als *Seitenfläche* von $\mathfrak{P}$ bezeichnen.

§ 5. Annäherung an einen konvexen Körper durch konvexe Polyeder.

Es sei wieder $\mathfrak{K}$ ein beliebiger konvexer Körper mit dem Nullpunkte o im Inneren; es sei $f(x, y, z)$ seine Distanzfunktion und $\frac{1}{R}$ das Minimum, $\frac{1}{r}$ das Maximum von $f(\alpha, \beta, \gamma)$ auf der Kugelfläche $\alpha^2 + \beta^2 + \gamma^2 = 1$.

Es sei δ eine beliebige positive Größe, und konstruieren wir irgend ein Netz $\mathfrak{N}$ von lauter kongruenten Würfeln mit der Kante δ und Seitenflächen parallel den Koordinatenebenen, die nur in den Seitenflächen aneinanderstoßen und den ganzen Raum lückenlos überdecken. Es seien U Würfel des Netzes vorhanden, welche überhaupt wenigstens einen Punkt des Körpers $\mathfrak{K}$ enthalten, und es sei $\mathfrak{U}$ der Bereich dieser Würfel, $\bar{\mathfrak{U}}$ der Bereich aller anderen Würfel des Netzes. Der ganze abgeschlossene Bereich $\bar{\mathfrak{U}}$ enthält keinen Punkt von $\mathfrak{K}$; die Bereiche $\mathfrak{U}$ und $\bar{\mathfrak{U}}$ aber haben ihre Begrenzung gemeinsam. Danach ist auf der Begrenzung von $\mathfrak{U}$ überall $f(x, y, z) > 1$ und enthält $\mathfrak{U}$ den Körper $\mathfrak{K}$ ganz im Inneren.

Denken wir uns die sämtlichen Punkte $p_1, p_2, \dots, p_n$ aufgesucht, die als Eckpunkte der Würfel in $\mathfrak{U}$ auftreten und konstruieren wir das

kleinste, alle diese Punkte enthaltende konvexe Polyeder $\mathfrak{P} = (p_1, p_2, \dots, p_n)$, so wird dieses Polyeder $\mathfrak{P}$ gewiß alle Würfel aus $\mathfrak{U}$, also den ganzen Bereich $\mathfrak{U}$ enthalten, also ebenfalls den Körper $\mathfrak{R}$ ganz im Inneren enthalten.

Andererseits gibt es in jedem Würfel, der zum Bereiche $\mathfrak{U}$ beiträgt, wenigstens einen Punkt x, y, z von $\mathfrak{R}$, also wofür $f(x, y, z) \leq 1$ ist, und dann gilt gemäß § 3, (3) und (4) in dem ganzen betreffenden Würfel von der Kante δ jedesmal $f(x, y, z) \leq 1 + \frac{\sqrt{3}}{r} \delta$. Insbesondere gilt daher diese Relation für alle Punkte $p_1, p_2, \dots, p_n$, und werden alle diese Punkte und mithin auch das ganze Polyeder $\mathfrak{P}$ vollständig in dem Körper $\left[1 + \frac{\sqrt{3}}{r} \delta\right] \mathfrak{R}$ enthalten sein, der durch Dilatation des Körpers $\mathfrak{R}$ vom Punkte o aus im Verhältnisse $1 + \frac{\sqrt{3}}{r} \delta : 1$ entsteht. Dilatieren wir dann die Körper $\mathfrak{P}$ und $\left[1 + \frac{\sqrt{3}}{r} \delta\right] \mathfrak{R}$ vom Punkte o aus in dem umgekehrten Verhältnisse $1 : 1 + \frac{\sqrt{3}}{r} \delta$, so erkennen wir, daß das Polyeder

$$\mathfrak{Q} = \frac{1}{1 + \frac{\sqrt{3}}{r} \delta} \mathfrak{P}$$

ganz im Körper $\mathfrak{R}$ enthalten ist. Wir kommen damit zu dem folgenden wichtigen Satze:

Ist $\mathfrak{R}$ ein beliebiger konvexer Körper mit dem Punkte o im Inneren, so kann man stets zwei einander in bezug auf den Punkt o homothetische (ähnliche und ähnlich gelegene) Polyeder $\mathfrak{P}$ und $\mathfrak{Q}$ konstruieren, so daß $\mathfrak{R}$ ganz in $\mathfrak{P}$, andererseits $\mathfrak{Q}$ ganz in $\mathfrak{R}$ liegt und dabei das Dilatationsverhältnis zur Erzeugung von $\mathfrak{P}$ aus $\mathfrak{Q}$ beliebig wenig größer als 1 ist.

§ 6. Stützebenen eines konvexen Körpers.

Wir knüpfen an das letzte Ergebnis noch eine wichtige Bemerkung. Nach der Definition einer Stützebene in § 4 haben wir unter einer *Stützebene an einen konvexen Körper* eine solche Ebene zu verstehen, welche wenigstens einen Punkt der Begrenzung des Körpers enthält, aber sein Inneres ganz auf einer Seite liegen hat. Wir können nun den Satz beweisen:

Durch jeden Punkt der Begrenzung eines konvexen Körpers geht wenigstens eine Stützebene an den Körper.

Durch § 4 ist dieser Satz bereits für Polyeder sichergestellt. Es sei nun $\mathfrak{R}$ ein beliebiger konvexer Körper mit o im Inneren, und denken wir uns für ihn wie vorhin das Polyeder $\mathfrak{P}$ konstruiert. Es sei p_0 mit den

Koordinaten x_0, y_0, z_0 ein beliebiger Punkt auf der Begrenzung von $\mathfrak{R}$. Der Strahl von $\mathfrak{o}$ durch $\mathfrak{p}_0$ treffe die Begrenzung des Polyeders $\mathfrak{P}$ im Punkte $\mathfrak{p}$ und es sei

$$(10) \quad ux + vy + wz \leq 1$$

die Bedingung einer Stützebene durch $\mathfrak{p}$ an $\mathfrak{P}$, also etwa einer Seitenfläche von $\mathfrak{P}$, welche durch den Punkt $\mathfrak{p}$ geht.

Da $\mathfrak{P}$ den Körper $\mathfrak{R}$ enthält, gilt dann (10) für jeden beliebigen Punkt x, y, z in $\mathfrak{R}$. Da $\mathfrak{P}$ den Körper $\mathfrak{R}$ und dieser die Kugel vom Radius r mit $\mathfrak{o}$ als Mittelpunkt ganz enthält, andererseits $\mathfrak{P}$ in

$$\left[1 + \frac{\sqrt{3}}{r} \delta\right] \mathfrak{R}$$

und letzterer Körper in der Kugel vom Radius $\left[1 + \frac{\sqrt{3}}{r} \delta\right] R$ mit $\mathfrak{o}$ als Mittelpunkt enthalten ist, wird

$$(11) \quad \frac{1}{r} \geq \sqrt{u^2 + v^2 + w^2} \geq \frac{1}{\left[1 + \frac{\sqrt{3}}{r} \delta\right] R}$$

sein. Insbesondere gilt (10) für den Punkt $\mathfrak{p}_0$; andererseits ist das Verhältnis der Längen $\mathfrak{o}\mathfrak{p} : \mathfrak{o}\mathfrak{p}_0 \leq 1 + \frac{\sqrt{3}}{r} \delta$ und $\frac{\mathfrak{p}_0\mathfrak{p}}{\mathfrak{o}\mathfrak{p}} \leq \frac{\sqrt{3}}{r} \delta$, danach folgt

$$(12) \quad 0 \leq 1 - (ux_0 + vy_0 + wz_0) \leq \frac{\sqrt{3}}{r} \delta.$$

Denken wir uns nun für δ eine unendliche Reihe nach der Grenze Null konvergierender positiver Werte angenommen, und bestimmen jedesmal in der hier dargelegten Weise zum Punkte $\mathfrak{p}_0$ ein System u, v, w , so werden die betreffenden unendlich vielen Systeme u, v, w nach der ersten Ungleichung in (11) wenigstens eine Häufungsstelle u_0, v_0, w_0 besitzen müssen; dabei werden wegen der zweiten Ungleichung in (12) diese Werte $u_0, v_0, w_0 \neq 0, 0, 0$ sein. Dann gilt wegen (10) für jeden Punkt x, y, z in $\mathfrak{R}$ die Ungleichung

$$(13) \quad u_0x + v_0y + w_0z \leq 1,$$

und wegen (13) wird für den Punkt x_0, y_0, z_0 in dieser Ungleichung das Gleichheitszeichen gelten. Also stellt (13) in der Tat die Bedingung einer durch $\mathfrak{p}_0$ gehenden Stützebene an $\mathfrak{R}$ dar.

Bedeutet $\mathfrak{q}$ einen Punkt außerhalb $\mathfrak{R}$, so trifft die Strecke $\mathfrak{o}\mathfrak{q}$ die Begrenzung von $\mathfrak{R}$ in einem Punkte $\mathfrak{p}_0$, und wir können durch $\mathfrak{p}_0$ eine Stützebene an $\mathfrak{R}$ legen. Eine parallele Ebene dazu durch einen, von $\mathfrak{p}_0$ und $\mathfrak{q}$ verschiedenen Punkt der Strecke $\mathfrak{p}_0\mathfrak{q}$ hat dann $\mathfrak{q}$ auf einer Seite, den Körper $\mathfrak{R}$ ganz auf der anderen Seite von sich liegen.

Wir geben für die hier entwickelten Sätze noch einen anderen Beweis.

Es sei q ein Punkt außerhalb $\mathfrak{R}$. Betrachten wir die Entfernung des Punktes q von einem in $\mathfrak{R}$ beliebig variablen Punkte x, y, z , so besitzt diese stetige Funktion von x, y, z , da $\mathfrak{R}$ eine abgeschlossene Punktmenge vorstellt, in $\mathfrak{R}$ ein bestimmtes Minimum. Es sei dann r ein Punkt aus $\mathfrak{R}$, für den dieses Minimum eintritt, so kann die Ebene durch r senkrecht zur Strecke qr keinen Punkt von $\mathfrak{R}$ auf derselben Seite wie q liegen haben. Denn wäre $\mathfrak{s}$ ein Punkt von $\mathfrak{R}$ auf derselben Seite dieser Ebene wie q , so würde mit r und $\mathfrak{s}$ die ganze Strecke $r\mathfrak{s}$ zu $\mathfrak{R}$ gehören; auf dieser Strecke aber befänden sich jedenfalls Punkte, deren Entfernung von q kleiner als die des Punktes r von q wären. Die Ebene durch q senkrecht auf qr enthält nunmehr keinen Punkt von $\mathfrak{R}$.

Jetzt sei p_0 ein Punkt der Begrenzung von $\mathfrak{R}$. Wir nehmen wieder an, daß $\mathfrak{R}$ den Nullpunkt o im Inneren enthalte. Wenn t ein Wert > 0 und < 1 ist, liegt dann der Körper $t\mathfrak{R}$ ganz im Inneren von $\mathfrak{R}$ und also p_0 stets außerhalb $t\mathfrak{R}$. Nach dem soeben Ausgeführten läßt sich infolgedessen eine Ebene durch p_0 legen, welche den Körper $t\mathfrak{R}$ ganz auf einer Seite läßt. Wir benutzen nun eine unendliche Reihe von wachsenden, nach der Grenze 1 konvergierenden Werten t , und stellen jedesmal in der angegebenen Weise dazu eine Ebene $ux + vy + wz = 1$ durch p_0 her; alsdann können wir ähnlich wie oben einsehen, daß die Wertsysteme u, v, w in den Gleichungen dieser Ebenen eine vom Systeme $0, 0, 0$ verschiedene Häufungsstelle u_0, v_0, w_0 besitzen, die uns eine Stützebene durch p_0 an $\mathfrak{R}$ bestimmt.

Die letzten Betrachtungen führen uns noch zu folgendem Satze:

Sind $\mathfrak{R}$ und $\mathfrak{R}^$ zwei verschiedene konvexe Körper, die keinen inneren Punkt miteinander gemein haben, so kann stets eine Ebene konstruiert werden, welche das Innere von $\mathfrak{R}$ ganz auf einer Seite, das Innere von $\mathfrak{R}^*$ ganz auf der anderen Seite liegen hat.*

Nehmen wir zuerst den Fall an, daß $\mathfrak{R}$ und $\mathfrak{R}^*$ überhaupt keinen Punkt gemein haben. In der Mannigfaltigkeit aller möglichen Wertsysteme von sechs reellen Variablen x, y, z, x^*, y^*, z^* bilden diejenigen dieser Systeme, bei denen x, y, z einen Punkt aus $\mathfrak{R}$ und x^*, y^*, z^* einen Punkt aus $\mathfrak{R}^*$ bedeuten, offenbar eine abgeschlossene Menge, weil $\mathfrak{R}$ sowohl wie $\mathfrak{R}^*$ für sich abgeschlossene Punktmengen sind. Infolgedessen gibt es unter den sämtlichen Entfernungen von je einem Punkte aus $\mathfrak{R}$ nach je einem Punkte aus $\mathfrak{R}^*$ ein bestimmtes Minimum und es seien z. B. r in $\mathfrak{R}$ und r^* in $\mathfrak{R}^*$ zwei Punkte, welche dieses Minimum der Entfernung zwischen $\mathfrak{R}$ und $\mathfrak{R}^*$ darbieten. Alsdann zeigt sich, daß die zur Strecke rr^* senkrechten Ebenen durch die Punkte dieser Strecke einen Parallelstreifen bilden, der das Innere von $\mathfrak{R}$ ganz auf einer Seite, das Innere von $\mathfrak{R}^*$ ganz auf der anderen Seite von sich liegen hat.

Haben jedoch $\mathfrak{R}$ und $\mathfrak{R}^*$ wenigstens einen Punkt ihrer Begrenzungen gemein, so wollen wir uns den Nullpunkt o im Inneren von $\mathfrak{R}$ gelegen denken, ersetzen $\mathfrak{R}$ durch $t\mathfrak{R}$, wobei $t > 0$ und < 1 ist, und gelangen zum Beweise unseres Satzes, indem wir eine Reihe nach der Grenze 1 wachsender Werte t heranziehen.

§ 7. Volumen eines konvexen Körpers. Schwerpunkt.

Für die Begründung des Volumenbegriffs ist folgende Bemerkung wichtig: Ein Bereich sei so beschaffen, daß er sich in eine endliche Anzahl von rechtwinkligen Parallelepipeda mit Seitenflächen parallel den Koordinatenebenen, also Bereichen von der Art

$$x_0 \leq x \leq x_0 + a, \quad y_0 \leq y \leq y_0 + b, \quad z_0 \leq z \leq z_0 + c$$

zerschneiden läßt, und enthalte einen zweiten in derselben Art zu zerlegenden Bereich. Indem man alle Ebenen, in welchen Seitenflächen der beiden Reihen von Parallelepipeda liegen, ganz durch die beiden Bereiche hindurchführt, erkennt man leicht, daß die Summe $\sum abc$ über alle Produkte abc für die Parallelepipeda des ersteren Bereichs stets größer ist als die entsprechende Summe in bezug auf den zweiten Bereich.

Es sei jetzt $\mathfrak{R}$ ein beliebiger konvexer Körper. Indem wir eine Parallelverschiebung der Koordinatenachsen zulassen, denken wir uns den Anfangspunkt o als inneren Punkt von $\mathfrak{R}$ und führen für $\mathfrak{R}$ die Funktion $f(x, y, z)$ und die Radien r und R wie in § 2 und § 3 ein. Es sei ε irgendeine positive GröÙe < 1 , und betrachten wir irgendein den Raum erfüllendes Netz $\mathfrak{N}$ von Würfeln mit einer Kante δ und gleichzeitig ein zweites solches Netz $\mathfrak{N}'$ von Würfeln mit einer Kante δ' , wobei δ wie δ' beide $< \frac{r}{\sqrt{3}}\varepsilon$ seien. Es bedeute U die Anzahl aller der Würfel des Netzes $\mathfrak{N}$, welche überhaupt wenigstens einen Punkt von $\mathfrak{R}$ enthalten, und A die Anzahl aller solchen Würfel aus $\mathfrak{N}$, welche mit allen ihren Punkten ganz ins Innere von $\mathfrak{R}$ fallen, ferner bedeute $\mathfrak{U}$ den Bereich jener U Würfel und falls $A > 0$ ist, $\mathfrak{A}$ den Bereich dieser A Würfel. Endlich mögen U' , A' , $\mathfrak{U}'$, $\mathfrak{A}'$ die entsprechenden Anzahlen und Bereiche in bezug auf das zweite Netz $\mathfrak{N}'$ vorstellen wie U , A , $\mathfrak{U}$, $\mathfrak{A}$ in bezug auf $\mathfrak{N}$.

Weil $\mathfrak{R}$ und daher auch $\mathfrak{A}$ ganz in dem Würfel $|x| \leq R$, $|y| \leq R$, $|z| \leq R$ liegt, hat man nach der am Eingange dieses Paragraphen gemachten Bemerkung

$$(14) \quad A\delta^3 \leq 8R^3.$$

Weil $\mathfrak{R}$ und daher auch $\mathfrak{U}'$ den Würfel $|x| \leq \frac{1}{\sqrt{3}}r$, $|y| \leq \frac{1}{\sqrt{3}}r$, $|z| \leq \frac{1}{\sqrt{3}}r$ ganz enthält, so gilt andererseits

$$(15) \quad U' \delta'^3 \geq \frac{8}{\sqrt{27}} r^3.$$

Da $\mathfrak{U}'$ den Körper $\mathfrak{R}$ und $\mathfrak{R}$ den Bereich $\mathfrak{A}$ ganz enthält, gilt

$$(16) \quad U' \delta'^3 - A \delta^3 \geq 0.$$

Ist $p(x, y, z)$ ein Punkt im Körper $(1 - \varepsilon)\mathfrak{R}$, so gilt dafür $f(x, y, z) \leq 1 - \varepsilon < 1 - \frac{\sqrt{3}}{r} \delta$; alsdann folgt für jeden solchen Punkt, welcher mit p sich in einem und dem nämlichen Würfel von $\mathfrak{R}$ befindet, wenigstens (s. § 6) $f(x, y, z) < 1$. Danach muß jedenfalls der Körper $(1 - \varepsilon)\mathfrak{R}$ mit allen seinen Punkten in den Bereich $\mathfrak{A}$ aufgenommen werden, und ist gewiß $A > 0$.

Jeder Würfel von $\mathfrak{U}'$ andererseits enthält wenigstens einen Punkt, wofür $f(x, y, z) \leq 1$ ist, und gilt dann in dem ganzen Würfel für alle Punkte $f(x, y, z) \leq 1 + \frac{\sqrt{3}}{r} \delta < 1 + \varepsilon$. Also liegt $\mathfrak{U}'$ ganz im Körper $(1 + \varepsilon)\mathfrak{R}$.

Dilatiert man nun den Bereich $\mathfrak{A}$, welcher $(1 - \varepsilon)\mathfrak{R}$ enthält, vom Nullpunkte o aus in allen Richtungen im Verhältnisse $1 + \varepsilon : 1 - \varepsilon$, so entsteht ein Bereich, der den Körper $(1 + \varepsilon)\mathfrak{R}$ und daher auch den Bereich $\mathfrak{U}'$ ganz in sich enthält. Danach ist auf Grund der am Anfange gemachten Bemerkung

$$\left(\frac{1 + \varepsilon}{1 - \varepsilon}\right)^3 A \delta^3 \geq U' \delta'^3.$$

Hieraus und aus (16) und (14) entnehmen wir nunmehr

$$(17) \quad 0 \leq U' \delta'^3 - A \delta^3 \leq 8R^3 \left(\left(\frac{1 + \varepsilon}{1 - \varepsilon}\right)^3 - 1 \right).$$

In dieser Relation können wir noch $U' \delta'^3$ mit $U \delta^3$ und andererseits $A \delta^3$ mit $A' \delta'^3$ vertauschen.

Diese Beziehung (17) zeigt uns, daß mit abnehmender Kante δ eines Würfelnetzes $\mathfrak{R}$ die daraus in bezug auf den Körper $\mathfrak{R}$ abgeleiteten Größen $A \delta^3$ und $U \delta^3$ nach einem bestimmten und beide nach dem nämlichen Grenzwerte V konvergieren; die Ungleichung (15) läßt noch erkennen, daß dieser Grenzwert V stets > 0 ist. Dieser Grenzwert heißt das *Volumen* des Körpers $\mathfrak{R}$.

Wir stellen sogleich noch die Existenz des *Schwerpunktes* eines konvexen Körpers fest. Bilden wir das Raumintegral

$$\iiint x \, dx \, dy \, dz,$$

zunächst erstreckt über die Bereiche $\mathcal{U}'$ und $\mathcal{A}$, so finden wir bestimmte Werte, die wir gleich

$$(18) \quad U' \delta'^3 x_{\mathcal{U}'} \quad \text{bzw.} \quad A \delta^3 x_{\mathcal{A}}$$

setzen. Nun enthält der Integrationsbereich $\mathcal{U}'$ ganz den Bereich $\mathcal{A}$, in $\mathcal{U}'$ gilt überdies stets $|x| < (1 + \varepsilon)R$. Danach ist der Betrag der Differenz dieser zwei Werte

$$|U' \delta'^3 x_{\mathcal{U}'} - A \delta^3 x_{\mathcal{A}}| \leq (1 + \varepsilon)R(U' \delta'^3 - A \delta^3).$$

Mit Rücksicht auf (17) geht daraus hervor, daß mit abnehmender Kante der Würfelnetze die Größen in (18) nach einer bestimmten Grenze konvergieren, die wir als das Raumintegral $\iiint x dx dy dz$ über $\mathfrak{K}$ bezeichnen und $= V x_{\mathfrak{K}}$ schreiben, worin V das Volumen von $\mathfrak{K}$ bedeute. Durch analoge Behandlung der y - und z -Koordinaten in $\mathfrak{K}$ gelangen wir zu zwei entsprechenden Grenzwerten $y_{\mathfrak{K}}$ und $z_{\mathfrak{K}}$. Der Punkt, dessen Koordinaten $x_{\mathfrak{K}}, y_{\mathfrak{K}}, z_{\mathfrak{K}}$ sind, erweist sich sodann als unveränderlich bei Einführung anderer Parallelkoordinaten an Stelle von x, y, z ; dieser Punkt heißt der *Schwerpunkt* des Körpers $\mathfrak{K}$.

Der Schwerpunkt eines konvexen Körpers ist stets ein innerer Punkt des Körpers.

Denn ist p irgendein Punkt außerhalb oder auf der Begrenzung von $\mathfrak{K}$, so können wir nach § 6 eine Ebene durch p legen, welche auf einer Seite keinen Punkt von $\mathfrak{K}$ liegen hat. Ist $Z = 0$ die Gleichung dieser Ebene und haben wir in $\mathfrak{K}$ stets $Z \geq 0$, so gilt in $\mathcal{A}$ stets $Z > 0$ und ist das Integral $\iiint Z dx dy dz$ über $\mathcal{A}$ positiv und gewiß nicht größer als das entsprechende Integral über $\mathfrak{K}$. Dadurch stellt sich der Wert Z für den Schwerpunkt von $\mathfrak{K}$ als > 0 heraus, es kann mithin dieser Punkt niemals mit p identisch sein.

II. Kapitel.

Die konvexen Körper als Ebenengebilde.

§ 8. Die Stützebenenfunktion eines konvexen Körpers mit dem Nullpunkte im Inneren.

Es sei zunächst $\mathfrak{K}$ wieder ein solcher konvexer Körper, welcher den Nullpunkt o als inneren Punkt enthält. Sind u, v, w irgendwelche feste Werte, die nicht sämtlich Null sind, und denkt man sich x, y, z als Koordinaten eines Punktes in $\mathfrak{K}$ variabel, so besitzt der Ausdruck

$$(19) \quad ux + vy + wz,$$

da $\mathfrak{K}$ eine abgeschlossene und ganz im Endlichen gelegene Punktmenge

vorstellt, in $\mathfrak{K}$ einen bestimmten größten Wert. Wir bezeichnen dieses *Maximum* von (19) in $\mathfrak{K}$ mit $H(u, v, w)$. Alsdann gilt für alle Punkte x, y, z in $\mathfrak{K}$ stets

$$(20) \quad ux + vy + wz \leq H(u, v, w),$$

und zugleich gibt es stets wenigstens einen solchen Punkt in $\mathfrak{K}$, für den in dieser Ungleichung das Gleichheitszeichen statthat, so daß (20) die Bedingung einer Stützebene an $\mathfrak{K}$ vorstellt. Da der Nullpunkt $\mathfrak{o}$ als innerer Punkt von $\mathfrak{K}$ nicht in dieser Ebene liegen kann, ist dann jedesmal $H(u, v, w) > 0$. — Für das System $u, v, w = 0, 0, 0$ verschwindet der Ausdruck (19) identisch und setzen wir demgemäß $H(0, 0, 0) = 0$.

Die hiermit definierte Funktion $H(u, v, w)$ nennen wir die *Stützebenenfunktion* des Körpers $\mathfrak{K}$. Diese Funktion ergibt in einfacher Weise die sämtlichen Stützebenen an den Körper $\mathfrak{K}$. Wir wollen eine nicht durch den Nullpunkt gehende Ebene, deren Gleichung

$$ux + vy + wz = 1$$

mit irgendwelchen Konstanten $u, v, w \neq 0, 0, 0$ lautet, kurz als die Ebene (u, v, w) bezeichnen. Zuzufolge (20) wird nun eine solche Ebene (u, v, w) den Körper $\mathfrak{K}$ nicht treffen und ihn ganz auf der Seite des Nullpunktes liegen haben, oder aber eine Stützebene an $\mathfrak{K}$ sein oder auch einen Teil von $\mathfrak{K}$ auf der dem Nullpunkte entgegengesetzten Seite liegen haben, je nachdem

$$H(u, v, w) < 1 \text{ oder } = 1 \text{ oder } > 1$$

ist. Die Gesamtheit aller Ebenen (u, v, w) , welche den Körper $\mathfrak{K}$ nicht zerschneiden, wird danach durch die Bedingung $H(u, v, w) \leq 1$ und die Gesamtheit aller Stützebenen an $\mathfrak{K}$ wird durch die Bedingung $H(u, v, w) = 1$ definiert.

Der Körper $\mathfrak{K}$ ist genau der Bereich aller derjenigen Punkte x, y, z , für welche bei jedem beliebigen Systeme u, v, w stets

$$(20) \quad ux + vy + wz \leq H(u, v, w)$$

gilt.

Denn für jeden Punkt x, y, z aus $\mathfrak{K}$ bestehen diese sämtlichen Ungleichungen. Ist aber $p(x, y, z)$ ein Punkt außerhalb $\mathfrak{K}$, so gibt es auf der Strecke vom Nullpunkte $\mathfrak{o}$ nach p einen Punkt p_0 der Begrenzung von $\mathfrak{K}$. Ist sodann (u_0, v_0, w_0) eine Stützebene durch p_0 an $\mathfrak{K}$, so gilt dabei $H(u_0, v_0, w_0) = 1$ und wird für p :

$$(21) \quad u_0x + v_0y + w_0z > 1 = H(u_0, v_0, w_0);$$

es bestehen also für p nicht sämtliche Ungleichungen (20).

Danach ist ein konvexer Körper $\mathfrak{K}$ mit dem Nullpunkte im Inneren durch Angabe seiner Stützebenenfunktion $H(u, v, w)$ jedesmal völlig bestimmt.

Die Funktion $H(u, v, w)$ besitzt die folgenden Eigenschaften:

1. Wenn $u, v, w \neq 0, 0, 0$ ist, gilt stets $H(u, v, w) > 0$. Ferner ist $H(0, 0, 0) = 0$.

2. Aus der Definition dieser Funktion ergibt sich unmittelbar:

$$(22) \quad H(tu, tv, tw) = tH(u, v, w), \text{ wenn } t > 0 \text{ ist.}$$

3. Für beliebige zwei Systeme $u_1, v_1, w_1; u_2, v_2, w_2$ gilt allgemein

$$(23) \quad H(u_1, v_1, w_1) + H(u_2, v_2, w_2) \geq H(u_1 + u_2, v_1 + v_2, w_1 + w_2).$$

Denn nach (20) gilt für jeden Punkt x, y, z in $\mathfrak{R}$

$$u_1x + v_1y + w_1z \leq H(u_1, v_1, w_1), \quad u_2x + v_2y + w_2z \leq H(u_2, v_2, w_2),$$

also auch

$$(u_1 + u_2)x + (v_1 + v_2)y + (w_1 + w_2)z \leq H(u_1, v_1, w_1) + H(u_2, v_2, w_2).$$

Es gibt aber andererseits wenigstens einen solchen Punkt x, y, z in $\mathfrak{R}$, für den die linke Seite der vorstehenden Ungleichung genau

$$= H(u_1 + u_2, v_1 + v_2, w_1 + w_2)$$

ausfällt, und dadurch folgt alsdann die Ungleichung (23).

Denken wir uns jetzt zu jeder Ebene (u, v, w) , welche den Körper $\mathfrak{R}$ nicht zerschneidet, den Punkt mit den Koordinaten u, v, w , also den *Pol* der Ebene in bezug auf die Kugelfläche $\mathfrak{G}$:

$$x^2 + y^2 + z^2 = 1$$

konstruiert. Aus den Sätzen in § 3 erkennen wir alsdann nach diesen Eigenschaften von $H(u, v, w)$, daß die Gesamtheit aller dieser Pole u, v, w , unter Hinzunahme noch des Nullpunktes, also aller Punkte u, v, w mit der Bedingung $H(u, v, w) \leq 1$ einen gewissen konvexen Körper $\mathfrak{S}$ mit dem Nullpunkt im Inneren erfüllt. Dieser Körper $\mathfrak{S}$ möge der *polare* Körper zu $\mathfrak{R}$ heißen.

Für einen beliebigen Punkt x, y, z aus $\mathfrak{R}$ und einen beliebigen Punkt u, v, w aus $\mathfrak{S}$ gilt nach (20) stets

$$(24) \quad ux + vy + wz \leq 1.$$

Der Körper $\mathfrak{S}$ ist dem Obigen zufolge genau der Bereich aller derjenigen einzelnen Systeme u, v, w , für welche bei *jedem* Punkte x, y, z aus $\mathfrak{R}$ stets diese Ungleichung (24) gilt, und wird eben durch diese Eigenschaft aus $\mathfrak{R}$ abgeleitet.

Andererseits ist nun $\mathfrak{R}$ genau der Bereich aller einzelnen Punkte x, y, z , für welche bei *jedem* Systeme u, v, w aus $\mathfrak{S}$ die Ungleichung (24) besteht. Denn, wie wir bei (21) sahen, gibt es, wenn x, y, z ein Punkt *außerhalb* $\mathfrak{R}$ ist, stets ein solches System u, v, w , wofür $H(u, v, w) = 1$ ist und (24) *nicht* gilt. Danach erkennen wir, daß die Beziehung der beiden konvexen Körper $\mathfrak{R}$ und $\mathfrak{S}$ hier eine reziproke ist. Der Körper $\mathfrak{R}$

ist wieder der polare Körper zu $\mathfrak{H}$. Zugleich leuchtet jetzt aus den Sätzen in § 3 ein, daß eine jede Funktion $H(u, v, w)$, welche die Bedingungen 1., 2., 3. erfüllt, stets die Stützebenenfunktion eines bestimmten konvexen Körpers $\mathfrak{K}$ mit dem Nullpunkte im Inneren ist.

Wir fügen noch folgende Bemerkungen über diesen Begriff des polaren Körpers hinzu. Bezeichnet $\mathfrak{H}$ den polaren Körper zu $\mathfrak{K}$ und ist t ein positiver Wert, so ist der polare Körper zu $t\mathfrak{K}$ der Körper $\frac{1}{t}\mathfrak{H}$. Enthält $\mathfrak{K}$ einen anderen konvexen Körper $\mathfrak{K}^*$ mit $\mathfrak{o}$ im Inneren, so enthält umgekehrt der polare Körper $\mathfrak{H}^*$ zu $\mathfrak{K}^*$ den polaren Körper $\mathfrak{H}$ zu $\mathfrak{K}$. Endlich zeigen wir, daß der polare Körper zu einem Polyeder mit $\mathfrak{o}$ im Inneren stets wieder ein Polyeder ist.

Bedeutet $\mathfrak{K}$ ein Polyeder, so läßt sich nach dem Satze in § 4 der Bereich der Punkte x, y, z in $\mathfrak{K}$ bereits vollständig durch eine endliche Anzahl von Ungleichungen

$$(25) \quad u^{(x)}x + v^{(x)}y + w^{(x)}z \leq 1 \quad (x = 1, 2, \dots, \nu)$$

charakterisieren. Sowie eine Ungleichung $\alpha u + \beta v + \gamma w \leq d$, wobei $\alpha^2 + \beta^2 + \gamma^2 = 1$ und $d > 0$ ist, für alle ν Systeme $u, v, w = u^{(x)}, v^{(x)}, w^{(x)}$ gilt, ist jedesmal $\frac{\alpha}{d}, \frac{\beta}{d}, \frac{\gamma}{d}$ ein Punkt in $\mathfrak{K}$. Es bedeute $\mathfrak{h}^{(x)}$ den Punkt mit den Koordinaten $u^{(x)}, v^{(x)}, w^{(x)}$ für $x = 1, 2, \dots, \nu$. Zudem ist noch $\mathfrak{K}$ als Polyeder ganz im Endlichen gelegen, d. h. in einer gewissen Kugel $x^2 + y^2 + z^2 \leq R^2$ enthalten. Danach kann es nun keine Ebene mit einer Distanz $< \frac{1}{R}$ von $\mathfrak{o}$ geben, welche alle die ν Punkte $\mathfrak{h}^{(x)}$ ganz auf einer Seite von sich liegen hat, d. h. der kleinste, alle diese Punkte enthaltende konvexe Bezirk $\mathfrak{H}^* = (\mathfrak{h}^{(1)}, \mathfrak{h}^{(2)}, \dots, \mathfrak{h}^{(\nu)})$ enthält notwendig den Nullpunkt im Inneren und ist also ein Polyeder. Bedeutet jetzt $\mathfrak{H}$ den polaren Körper zu $\mathfrak{K}$, so enthält $\mathfrak{H}$ jeden der Punkte $\mathfrak{h}^{(x)}$ und also das ganze Polyeder $\mathfrak{H}^*$. Ist dann $\mathfrak{K}^*$ der zu $\mathfrak{H}^*$ polare Körper, so erfüllt jeder Punkt x, y, z dieses Körpers alle ν Ungleichungen (25), also ist $\mathfrak{K}^*$ ganz in $\mathfrak{K}$ enthalten und enthält daher umgekehrt $\mathfrak{H}^*$ den Körper $\mathfrak{H}$. Danach ist der polare Körper zu $\mathfrak{K}$ identisch mit dem Polyeder $(\mathfrak{h}^{(1)}, \mathfrak{h}^{(2)}, \dots, \mathfrak{h}^{(\nu)})$.

§ 9. Kugeln, die in einem konvexen Körper enthalten sind. Homothetische Körper.

Zu jeder Richtung (α, β, γ) haben wir in

$$(26) \quad \alpha x + \beta y + \gamma z \leq H(\alpha, \beta, \gamma)$$

eine Stützebene an $\mathfrak{K}$ mit dieser Richtung als (äußerer) Normale, so daß also auf der Seite dieser Ebene, nach der jene Richtung weist, kein Punkt von $\mathfrak{K}$ liegt. Der senkrechte Abstand dieser Stützebene vom Nullpunkte ist $H(\alpha, \beta, \gamma)$.

Bilden wir gemäß § 2 die Distanzfunktion $f(x, y, z)$ des Körpers $\mathfrak{R}$ und wenden die Ungleichung (26) auf den Punkt der Begrenzung von $\mathfrak{R}$ an, der in der Richtung (α, β, γ) von o aus liegt, so folgt

$$\frac{1}{f(\alpha, \beta, \gamma)} \leq H(\alpha, \beta, \gamma).$$

Es sei $\mathfrak{R}^*$ ein zweiter konvexer Körper mit o im Inneren und $H^*(u, v, w)$ die Stützebenenfunktion von $\mathfrak{R}^*$. Ist der Körper $\mathfrak{R}^*$ ganz in $\mathfrak{R}$ enthalten, so gilt für jede Richtung (α, β, γ) stets

$$(27) \quad H^*(\alpha, \beta, \gamma) \leq H(\alpha, \beta, \gamma),$$

und wegen der Regel 2. in § 8 dann überhaupt allgemein $H^*(u, v, w) \leq H(u, v, w)$. Umgekehrt, wenn allgemein die Ungleichungen (27) statt haben, so ist nach (20) der Körper $\mathfrak{R}^*$ ganz im Körper $\mathfrak{R}$ enthalten.

Ist insbesondere $\mathfrak{R}$ ein Polyeder und gilt die Ungleichung (27) für diejenigen Richtungen (α, β, γ) , welche die äußeren Normalen der Seitenflächen von $\mathfrak{R}$ abgeben, so ist nach § 4 bereits $\mathfrak{R}^*$ in $\mathfrak{R}$ enthalten und gilt also jene Ungleichung bereits für jede mögliche Richtung.

Für die Kugel $x^2 + y^2 + z^2 \leq t^2$ ($t > 0$) ist $H^*(\alpha, \beta, \gamma)$ konstant gleich dem Radius t . Es sei wie in § 3 r das Minimum, R das Maximum von $\frac{1}{f(\alpha, \beta, \gamma)}$ auf der Kugeloberfläche $\mathfrak{G}$, also r der Radius der größten in $\mathfrak{R}$ enthaltenen, R der Radius der kleinsten, $\mathfrak{R}$ enthaltenden Kugel mit dem Nullpunkte als Mittelpunkt. Alsdann gilt nach (27) für $\mathfrak{R}$ stets $r \leq H(\alpha, \beta, \gamma)$ und $H(\alpha, \beta, \gamma) \leq R$, aber, wenn t ein Wert $> r$ und $< R$ ist, weder stets $t \leq H(\alpha, \beta, \gamma)$ noch stets $H(\alpha, \beta, \gamma) \leq t$. Also bedeuten zugleich r das Minimum und R das Maximum von $H(\alpha, \beta, \gamma)$ auf der ganzen Kugeloberfläche $\alpha^2 + \beta^2 + \gamma^2 = 1$.

Sind a, b, c irgendwelche feste Werte, so stellt der Ausdruck

$$(28) \quad H(\alpha, \beta, \gamma) - a\alpha - b\beta - c\gamma$$

den Abstand des Punktes a, b, c von der Stützebene an $\mathfrak{R}$ mit der Normale α, β, γ vor, und zwar, falls die Stützebene den Körper und den Punkt voneinander trennt, diesen Abstand noch mit negativem Vorzeichen versehen. Der Ausdruck (28) besitzt auf der Kugeloberfläche $\mathfrak{G}$ ein bestimmtes Minimum, das wir mit $r(a, b, c)$ bezeichnen. Dieses Minimum ist immer dann und nur dann noch positiv, wenn a, b, c ein innerer Punkt von $\mathfrak{R}$ ist, und in solchem Falle bedeutet $r(a, b, c)$ den Radius der größten Kugel mit a, b, c als Mittelpunkt, welche noch ganz im Körper $\mathfrak{R}$ enthalten ist. Weiter stellt der Wert $r(a, b, c)$ offenbar eine stetige Funktion der Argumente a, b, c vor und besitzt deshalb in dem abgeschlossenen Bereiche der Punkte des Körpers $\mathfrak{R}$ ein bestimmtes Maximum, das $r_{\max}$ heiße und alsdann zugleich das Maximum unter den Radien

aller ganz in $\mathfrak{K}$ enthaltenen Kugeln bedeutet. Insbesondere ist daher $r(0, 0, 0) \leq r_{\max}$. — Wir werden in der Folge (§ 23) beweisen, daß, wenn a, b, c speziell die Koordinaten des *Schwerpunktes* des Körpers $\mathfrak{K}$ bedeuten, jedenfalls $r(a, b, c) \geq \frac{1}{2} r_{\max}$ ist.

Die zu einer Richtung (α, β, γ) entgegengesetzte Richtung hat als Kosinus ihrer Neigungswinkel gegen die Achsen $-\alpha, -\beta, -\gamma$. Die Stützebene an $\mathfrak{K}$ mit dieser Richtung $(-\alpha, -\beta, -\gamma)$ als (äußerer) Normale hat vom Nullpunkte den Abstand $H(-\alpha, -\beta, -\gamma)$ und $\mathfrak{K}$ befindet sich sodann ganz in dem von diesen zwei parallelen Stützebenen begrenzten Streifen

$$(29) \quad -H(-\alpha, -\beta, -\gamma) \leq \alpha x + \beta y + \gamma z \leq H(\alpha, \beta, \gamma).$$

Der gegenseitige Abstand dieser zwei parallelen Stützebenen ist $H(\alpha, \beta, \gamma) + H(-\alpha, -\beta, -\gamma)$. Diese Größe heiße die *Breite* des Körpers $\mathfrak{K}$ in der Richtung α, β, γ (oder $-\alpha, -\beta, -\gamma$); sie ist jedenfalls stets $\geq 2r_{\max}$.

Wir entwickeln hier noch ein in der Folge nützliches Kriterium zur Entscheidung, ob zwei gegebene konvexe Körper $\mathfrak{K}$ und $\mathfrak{K}'$ *homothetisch* sind (auseinander durch *Translation* und *Dilatation* hervorgehen) oder nicht. Wir können jeden der Körper durch eine bestimmte Translation so verschieben, daß sein Schwerpunkt in den Nullpunkt zu liegen kommt. Denken wir uns der Einfachheit wegen, daß solches für beide Körper von vornherein der Fall ist; sind die Körper homothetisch, so muß dann der eine Körper aus dem anderen durch eine Dilatation von dem gemeinsamen Schwerpunkte abzuleiten sein. Es seien nun $H(u, v, w)$ und $H'(u, v, w)$ die Stützebenenfunktionen, ferner V und V' die Volumina von $\mathfrak{K}$ und $\mathfrak{K}'$. Der Körper $\mathfrak{K}^* = \sqrt[3]{\frac{V}{V'}} \mathfrak{K}'$ hat dann dasselbe Volumen V wie $\mathfrak{K}$ und muß daher, wenn $\mathfrak{K}$ und $\mathfrak{K}'$ homothetisch sind, sich mit $\mathfrak{K}$ decken. Fallen jedoch $\mathfrak{K}^*$ und $\mathfrak{K}$ nicht zusammen, so kann auch nicht $\mathfrak{K}^*$ ganz in $\mathfrak{K}$ enthalten sein, da $\mathfrak{K}^*$ sonst ein kleineres Volumen wie $\mathfrak{K}$ besäße. Also muß es dann nach (27) wenigstens eine Richtung (α, β, γ) geben, wofür $\sqrt[3]{\frac{V}{V'}} H'(\alpha, \beta, \gamma) > H(\alpha, \beta, \gamma)$ ist.

Betrachten wir nun die Funktion

$$(30) \quad \sqrt[3]{\frac{V}{V'}} \frac{H'(\alpha, \beta, \gamma)}{H(\alpha, \beta, \gamma)},$$

deren Nenner niemals verschwindet, so hat diese Funktion auf der Kugel-
fläche $\mathfrak{E}$ ein bestimmtes *Maximum*, das Q heiße. Alsdann gilt $Q = 1$
oder $Q > 1$, je nachdem $\mathfrak{K}$ und $\mathfrak{K}'$ homothetisch sind oder nicht.

Die Größe Q bedeutet zugleich den kleinsten Wert, wofür noch alle
Ungleichungen

$$\sqrt[3]{\frac{V}{V'}} H'(\alpha, \beta, \gamma) \leq Q H(\alpha, \beta, \gamma)$$

bestehen und also der Körper $\mathfrak{R}^*$ noch im Körper $Q\mathfrak{R}$ enthalten ist. Für den Fall, daß $\mathfrak{R}$ ein Polyeder vorstellt, gilt nach einer oben gemachten Bemerkung diese Ungleichung für jede Richtung (α, β, γ) , sowie sie für die (äußeren) Normalen der Seitenflächen von $\mathfrak{R}$ erfüllt ist, und wird daher der Maximumwert Q von (30) bereits bei einer dieser letzteren Richtungen, die in endlicher Anzahl vorhanden sind, zum Vorschein kommen.

Andererseits besitzt die Funktion (30) auf der Kugeloberfläche $\mathfrak{E}$ ein bestimmtes, noch positives *Minimum*, das q heiße, und wird dann $q = 1$ oder $q < 1$ sein, je nachdem $\mathfrak{R}$ und $\mathfrak{R}'$ homothetisch sind oder nicht. Die Größen $Q - 1$ und $1 - q$ sind also beide gleichzeitig $= 0$ oder > 0 .

§ 10. Konvexe Körper in beliebiger Lage zum Nullpunkte.

Wir definieren jetzt die Stützebenenfunktion für einen konvexen Körper ohne Rücksicht auf dessen Lage zum Nullpunkte. Es sei $\mathfrak{R}$ ein beliebiger konvexer Körper und e mit den Koordinaten $x = a, y = b, z = c$ ein beliebiger angenommener Punkt im Inneren von $\mathfrak{R}$. Durch die Translation des Körpers $\mathfrak{R}$ vom Punkte a, b, c nach dem Nullpunkte o erhalten wir einen konvexen Körper $\mathfrak{R}^*$ mit o im Inneren; dieser besteht in der Menge der Punkte $x - a, y - b, z - c$, während x, y, z sich in $\mathfrak{R}$ bewegt. Es sei $H^*(u, v, w)$ die Stützebenenfunktion von $\mathfrak{R}^*$, so bedeutet $H^*(u, v, w)$ das Maximum von

$$u(x - a) + v(y - b) + w(z - c)$$

für die Punkte x, y, z in $\mathfrak{R}$. Das Maximum von $ux + vy + wz$ in $\mathfrak{R}$ wird dann durch den Ausdruck

$$(31) \quad H(u, v, w) = H^*(u, v, w) + au + bv + cw$$

dargestellt sein; dieses Maximum $H(u, v, w)$ bezeichnen wir wieder als die *Stützebenenfunktion* von $\mathfrak{R}$. Dabei erweist sich wieder $\mathfrak{R}$ genau als der Bereich derjenigen Punkte x, y, z , welche sämtlichen Ungleichungen

$$(32) \quad ux + vy + wz \leq H(u, v, w)$$

für alle möglichen Werte u, v, w genügen, so daß die Stützebenenfunktion jedenfalls wieder den konvexen Körper $\mathfrak{R}$ vollständig bestimmt.

Nach der Regel 1. in § 8 war stets $H^*(u, v, w) > 0$, wenn $u, v, w \neq 0, 0, 0$ ist. Diese Ungleichung braucht nun nicht mehr stets für $H(u, v, w)$ zu gelten; wir finden jedoch aus (31):

$$H(u, v, w) + H(-u, -v, -w) = H^*(u, v, w) + H^*(-u, -v, -w)$$

und können danach folgende Eigenschaften behaupten:

1. Für die Funktion $H(u, v, w)$ gilt stets

$$(33) \quad H(u, v, w) + H(-u, -v, -w) > 0, \text{ wenn } u, v, w \neq 0, 0, 0$$

ist. Außerdem ist $H(0, 0, 0) = 0$.

Aus den Regeln 2. und 3. für $H^*(u, v, w)$ gemäß § 8, finden wir genau wie früher:

$$2. \quad H(tu, tv, tw) = tH(u, v, w), \text{ wenn } t > 0 \text{ ist.}$$

$$3. \quad H(u_1, v_1, w_1) + H(u_2, v_2, w_2) \geq H(u_1 + u_2, v_1 + v_2, w_1 + w_2).$$

Wir beweisen jetzt umgekehrt den folgenden Satz:

Genügt eine Funktion $H(u, v, w)$ den vorstehenden Bedingungen 1., 2., 3., so ist sie jedesmal die Stützebenenfunktion eines bestimmten konvexen Körpers.

Wir haben zunächst festzustellen, daß unter der gemachten Voraussetzung stets drei Konstanten a, b, c auf solche Weise bestimmt werden können, daß die Funktion

$$H(u, v, w) - (au + bv + cw) = H^*(u, v, w)$$

für alle von $0, 0, 0$ verschiedenen Systeme u, v, w stets > 0 ausfällt. Alsdann wird nach den Resultaten in § 8 diese Funktion $H^*(u, v, w)$ die Stützebenenfunktion eines gewissen konvexen Körpers $\mathfrak{K}^*$ mit dem Nullpunkte im Inneren bilden, und $H(u, v, w)$ wird die Stützebenenfunktion desjenigen konvexen Körpers $\mathfrak{K}$ sein, der aus $\mathfrak{K}^*$ durch die Translation vom Nullpunkte nach dem Punkte a, b, c entsteht, dabei also den letzteren Punkt als inneren aufweist.

Gilt stets $H(u, v, w) > 0$, wenn $u, v, w \neq 0, 0, 0$ sind, so brauchen wir nur $a = 0, b = 0, c = 0$ zu nehmen.

Wir setzen jetzt voraus, es sei nicht für jedes System $u, v, w \neq 0, 0, 0$ stets $H(u, v, w) > 0$, wir nehmen aber zunächst an, daß wenigstens stets $H(u, v, 0) > 0$ gelte, so oft $u, v \neq 0, 0$ sind. Um den Gang des alsdann zu führenden Beweises vorausszusehen, bedenken wir, daß, wenn tatsächlich der durch die Ungleichungen (32) definierte Bereich ein konvexer Körper ist, die letzte Annahme darauf hinauskommen wird, daß die orthogonale Projektion dieses Körpers auf die xy -Ebene den Nullpunkt $x = 0, y = 0$ als inwendigen Punkt enthält. Alsdann steht zu erwarten, daß die z -Achse den Körper in einer Strecke $c'c''$ durchsetzt und jeder von den Endpunkten verschiedene Punkt dieser Strecke ein innerer Punkt des Körpers wird. Wir haben nun allein aus den Bedingungen 1., 2., 3. für die Funktion $H(u, v, w)$ abzuleiten, daß auf der z -Achse innere Punkte des Bereichs (32) vorhanden sind.

Wir setzen hiernach $a = 0, b = 0$ und haben alsdann eine Konstante c derart zu suchen, daß stets $H(u, v, w) - cw > 0$ ist, wenn $w \neq 0$ ist. Er-

setzen wir, wenn $w \neq 0$ ist, u, v durch uw, vw und beachten die Regel 1., so sollen also nach Voraussetzung *nicht* alle Ungleichungen $H(u, v, 1) > 0$, $H(-u, -v, -1) > 0$ gelten, es existiere etwa ein System $-u^{(0)}, -v^{(0)}, -1$, wofür $H(-u^{(0)}, -v^{(0)}, -1) = -c^{(0)} \leq 0$ ist. Verlangt wird, die Konstante c derart zu bestimmen, daß stets

$$(34) \quad H(u, v, 1) - c > 0, \quad H(-u, -v, -1) + c > 0$$

ist.

Für die Funktion $H(u, v, 0)$ der zwei Argumente u, v haben wir $H(u, v, 0) > 0$, wenn $u, v \neq 0, 0$ sind, $H(0, 0, 0) = 0$, ferner entnehmen wir aus 2. und 3.:

$$H(tu, tv, 0) = tH(u, v, 0), \text{ wenn } t > 0 \text{ ist,}$$

$$H(u_1, v_1, 0) + H(u_2, v_2, 0) \geq H(u_1 + u_2, v_1 + v_2, 0).$$

Durch entsprechende Überlegungen, wie sie in § 3 für die Funktion $f(x, y, z)$ angestellt wurden, erkennen wir hieraus, daß $H(u, v, 0)$ eine stetige Funktion von u, v ist, und schließen wir auf die Existenz zweier positiver Größen s und S derart, daß stets

$$(35) \quad \frac{1}{s} \sqrt{u^2 + v^2} \geq H(u, v, 0) \geq \frac{1}{S} \sqrt{u^2 + v^2}$$

ist. Die Beziehungen

$$-H(-u_2, -v_2, 0) \leq H(u_1 + u_2, v_1 + v_2, 1) - H(u_1, v_1, 1) \leq H(u_2, v_2, 0)$$

und die erste Ungleichung in (35) zeigen uns, daß $H(u, v, 1)$ eine stetige Funktion der Argumente u, v ist, und analog erkennen wir $H(-u, -v, -1)$ als stetige Funktion von u, v .

Aus

$$(36) \quad H(u, v, 1) + H(-u^{(0)}, -v^{(0)}, -1) \geq H(u - u^{(0)}, v - v^{(0)}, 0)$$

folgt, wenn $u, v \neq u^{(0)}, v^{(0)}$ ist, $H(u, v, 1) > c^{(0)}$. Für das System $u = u^{(0)}, v = v^{(0)}$ geht aus $H(u^{(0)}, v^{(0)}, 1) + H(-u^{(0)}, -v^{(0)}, -1) > 0$ nach Regel 1. die entsprechende Ungleichung hervor, so daß allgemein für jedes mögliche System u, v stets

$$(37) \quad H(u, v, 1) > c^{(0)}$$

gilt. Aus (36) unter Berücksichtigung der zweiten Ungleichung in (35) ergibt sich ferner, daß gewiß nur dann, wenn

$$(38) \quad \sqrt{(u - u^{(0)})^2 + (v - v^{(0)})^2} \leq S(H(u^{(0)}, v^{(0)}, 1) + H(-u^{(0)}, -v^{(0)}, -1))$$

gilt, der Wert $H(u, v, 1)$ sich $\leq H(u^{(0)}, v^{(0)}, 1)$ erweisen kann. In dem durch diese Ungleichung (38) definierten abgeschlossenen Bereiche nun besitzt $H(u, v, 1)$ ein bestimmtes Minimum, dessen Wert c'' sei und das etwa für $u = u'', v = v''$ eintrete. Für jedes beliebige System u, v gilt dann

$$(39) \quad H(u, v, 1) \geq c''.$$

Aus (37) folgt für $u = u'', v = v''$ insbesondere $c'' > c^{(0)}$.

Betrachten wir ferner die Funktion $H(-u, -v, -1)$; diese besitzt unter anderem den Wert $-c^{(0)}$, der ≤ 0 ist. Wegen

$$H(-u, -v, -1) + H(u^{(0)}, v^{(0)}, 1) \geq H(-u + u^{(0)}, -v + v^{(0)}, 0)$$

kann sie Werte, die $\leq H(-u^{(0)}, -v^{(0)}, -1)$ sind, jedenfalls nur wieder in dem durch (38) definierten Bereiche von Systemen u, v annehmen. In diesem abgeschlossenen Bereiche wird sie ein bestimmtes Minimum $-c'$ besitzen, wobei $c' \geq c^{(0)} \geq 0$ ist, und es sei $u = u', v = v'$ ein solches System, für das $H(-u', -v', -1) = -c'$ ist. Für jedes mögliche System u, v gilt dann

$$(40) \quad H(-u, -v, -1) \geq -c'.$$

Wir könnten nun an Stelle des Systems $u^{(0)}, v^{(0)}$ oben auch das System u', v' verwenden, und genau wie $c'' > c^{(0)}$ vorhin finden wir dann $c'' > c'$. Ist jetzt c irgendeine Konstante $> c'$ und $< c''$, so haben wir nach (39) und (40)

$$H(u, v, 1) > c \quad \text{und} \quad H(-u, -v, -1) > -c$$

und diese Konstante c entspricht daher den oben gestellten Anforderungen. —

Wir nehmen jetzt zweitens an, daß auch nicht durchweg $H(u, v, 0) > 0$ sei, wenn $u, v \neq 0, 0$ sind, daß aber wenigstens stets $H(u, 0, 0) > 0$ für ein $u \neq 0$ gelte, d. h. also daß $H(1, 0, 0)$ und $H(-1, 0, 0)$ beide > 0 seien. Alsdann können wir $a = 0$ setzen und durch eine entsprechende Überlegung für die Funktion $H(u, v, 0)$ von zwei Argumenten, wie wir sie soeben in bezug auf $H(u, v, w)$ anstellten, ermitteln wir eine Konstante b derart, daß $H(u, v, 0) - bv$ stets > 0 ist, wenn $v \neq 0$ ist. Damit bekommt die letztere Differenz den Charakter, den wir vorhin für $H(u, v, w)$ voraussetzten und wir können weiter c so wählen, daß noch

$$(H(u, v, w) - bv) - cw$$

stets > 0 ist, wenn $w \neq 0$ ist.

Sind endlich auch nicht $H(1, 0, 0)$ und $H(-1, 0, 0)$ beide positiv, so können wir a so wählen, daß $a < H(1, 0, 0)$ und $> -H(-1, 0, 0)$ ist. Damit wird $H(u, 0, 0) - au > 0$ für jedes $u \neq 0$, und wir können darauf wie soeben b und c nacheinander so bestimmen, daß

$$H(u, v, 0) - au - bv > 0$$

ist, sowie $v \neq 0$ ist, und endlich $H(u, v, w) - au - bv - cw > 0$, sowie $w \neq 0$ ist.

§ 11. Ovale. Konvexe Bezirke.

Eine Punktmenge, welche ganz in einer Ebene gelegen ist und welche dabei 1° mit einer beliebigen Geraden der Ebene sei es eine Strecke, sei

es einen Punkt, sei es keinen Punkt gemein hat, 2^o wenigstens drei nicht in einer Geraden gelegene Punkte aufweist, soll ein *Oval* heißen.*)

Wir können die Betrachtungen in §§ 1—10 über gewisse räumliche Gebilde sinngemäß auf die ebenen Ovale übertragen. Insbesondere zeigt sich (s. § 3), daß ein Oval stets ganz im Endlichen liegt, d. h. für die Koordinaten der sämtlichen Punkte in ihm obere und untere Grenzen existieren, stets abgeschlossen ist, ferner daß in der Ebene des Ovals durch jeden Punkt seiner *Berandung* (Begrenzung in der Mannigfaltigkeit der Ebene) wenigstens eine Gerade geht, welche nicht zu beiden Seiten von sich Punkte des Ovals liegen hat und die wir danach als eine *Stützgerade* an das Oval bezeichnen (s. § 6).

Eine Punktmenge, von welcher allein die Eigenschaft 1^o der konvexen Körper feststeht, daß sie mit einer jeden Geraden sei es eine Strecke, sei es einen Punkt, sei es keinen Punkt gemein hat, bezeichnen wir schlechthin als einen konvexen *Bezirk*. Ein beliebiger konvexer Bezirk bedeutet entweder einen konvexen Körper oder ein Oval in einer Ebene oder eine geradlinige Strecke oder einen einzelnen Punkt. Jene Grundeigenschaft 1^o der konvexen Bezirke ist nach § 3 gleichbedeutend mit dem Inbegriff der drei Eigenschaften:

1a) die Menge enthält mit irgend zwei Punkten stets auch die ganze dieselben verbindende Strecke;

1b) die Menge ist ganz im Endlichen gelegen;

1c) die Menge ist abgeschlossen.

Wir können jedem konvexen Bezirke $\mathfrak{R}$ eine *Stützebenenfunktion* $H(u, v, w)$ zuweisen; dabei definieren wir $H(u, v, w)$ bei beliebigen Argumenten u, v, w als das Maximum des Ausdrucks $ux + vy + wz$ in der Menge der Punkte x, y, z des Bezirks $\mathfrak{R}$. Dieser Bezirk $\mathfrak{R}$ ist sodann jedesmal durch die sämtlichen Ungleichungen

$$(41) \quad ux + vy + wz \leq H(u, v, w)$$

völlig charakterisiert. Bei einem beliebigen konvexen Bereich $\mathfrak{R}$ haben wir für die Stützebenenfunktion $H(u, v, w)$ im allgemeinen die Regeln 1., 2., 3. wie in § 10 für einen konvexen Körper, nur müssen wir in der unter (33) aufgeführten Ungleichung noch das Zeichen $=$ zulassen, wir können also nur die Bedingung

$$(42) \quad H(u, v, w) + H(-u, -v, -w) \geq 0,$$

auch wenn $u, v, w \neq 0, 0, 0$ sind, behaupten. Diese Ungleichung (42) geht daraus hervor, daß das Minimum von $ux + vy + wz$ in $\mathfrak{R}$ gleich $-H(-u, -v, -w)$ ist. Andererseits gilt nun der Satz:

*) Diese Bezeichnung ist dem Aufsätze von Herrn Brunn „Über Ovale und Eiflächen“ (Inaugural-Dissertation, München, 1887) entnommen.

Eine beliebige Funktion $H(u, v, w)$, welche allgemein die Bedingungen

1. $H(u, v, w) + H(-u, -v, -w) \geq 0, \quad H(0, 0, 0) = 0,$
2. $H(tu, tv, tw) = tH(u, v, w),$ wenn $t > 0$ ist,
3. $H(u_1, v_1, w_1) + H(u_2, v_2, w_2) \geq H(u_1 + u_2, v_1 + v_2, w_1 + w_2)$

erfüllt, ist jedesmal die Stützebenenfunktion eines bestimmten konvexen Bezirks.

In der Tat, gilt dabei in (42) stets das Zeichen $>$, wenn $u, v, w \neq 0, 0, 0$ sind, so ist nach § 10 dann $H(u, v, w)$ die Stützebenenfunktion eines bestimmten konvexen Körpers. Finden wir jedoch irgendein von $0, 0, 0$ verschiedenes System u^*, v^*, w^* , wofür die Gleichung

$$H(u^*, v^*, w^*) + H(-u^*, -v^*, -w^*) = 0$$

besteht, so zeigen die unter den Ungleichungen (41) enthaltenen zwei Bedingungen

$$-H(-u^*, -v^*, -w^*) \leq u^*x + v^*y + w^*z \leq H(u^*, v^*, w^*),$$

daß der durch diese Ungleichungen (41) definierte Bereich jedenfalls ganz in der Ebene

$$u^*x + v^*y + w^*z = H(u^*, v^*, w^*)$$

liegen muß. Wir denken uns nun eine solche Transformation der orthogonalen Koordinaten vorgenommen, daß in den neuen Koordinaten, für die wir wieder die alten Zeichen verwenden, diese Ebene die xy -Ebene wird, oder also, wir setzen $H(0, 0, 1) = -H(0, 0, -1) = 0$ voraus. Als dann haben wir

$$0 = -H(0, 0, -w) \leq H(u, v, w) - H(u, v, 0) \leq H(0, 0, w) = 0,$$

mithin wird allgemein $H(u, v, w) = H(u, v, 0)$ und das System der Ungleichungen (41) kommt auf das System der Bedingungen

$$(43) \quad z = 0, \quad ux + vy \leq H(u, v, 0)$$

hinaus. Ist sodann stets $H(u, v, 0) + H(-u, -v, 0) > 0$, wenn $u, v \neq 0, 0$ sind, so zeigen entsprechende Überlegungen wie in § 10, daß der hierdurch definierte Bereich ein Oval in der Ebene $z=0$ mit der Stützebenenfunktion $H(u, v, w) = H(u, v, 0)$ ist.

Finden wir weiter noch ein System $u, v \neq 0, 0$, wofür die eben erwähnte Summe $= 0$ ist, so erkennen wir dagegen, daß der Bereich (43) ganz in einer Geraden der xy -Ebene liegt. Unter Zulassung einer Koordinatentransformation in der xy -Ebene nehmen wir an, es sei

$$H(0, 1, 0) = -H(0, -1, 0) = 0.$$

Dann folgt allgemein $H(u, v, 0) = H(u, 0, 0)$ und muß der durch (43) definierte Bereich in die x -Achse fallen. Er stellt eine Strecke auf der

x -Achse vor, wenn $H(1, 0, 0) + H(-1, 0, 0) > 0$ ist. Wenn endlich auch die letztere Summe $= 0$ ist, gilt durchgehends $H(u, v, w) + H(-u, -v, -w) = 0$ und reduziert sich der Bereich (43) auf einen einzigen Punkt.

Aus der Bedeutung der Stützebenenfunktion eines konvexen Bezirks entnehmen wir sofort die Tatsache:

Sind $\mathfrak{R}$ und $\mathfrak{R}^$ zwei konvexe Bezirke mit den Stützebenenfunktionen H und H^* , so ist dann und nur dann $\mathfrak{R}$ in $\mathfrak{R}^*$ enthalten, wenn für jedes Wertsystem u, v, w stets*

$$(44) \quad H(u, v, w) \leq H^*(u, v, w)$$

gilt.

Wir erwähnen hier noch eine Methode, die dazu dient, mittels mehrerer konvexer Bezirke einen bestimmten weiteren solchen Bezirk zu definieren.

Sind $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ eine endliche Reihe von konvexen Bezirken mit den Stützebenenfunktionen $H_1, H_2, \dots, H_m$ und bedeutet $H(u, v, w)$ bei beliebigen Argumenten u, v, w jedesmal das Maximum unter den Werten

$$(45) \quad H_1(u, v, w), H_2(u, v, w), \dots, H_m(u, v, w),$$

so bildet $H(u, v, w)$ stets wieder die Stützebenenfunktion eines gewissen konvexen Bezirks. Dieser neue Bezirk, den wir mit $(\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m)$ bezeichnen wollen, enthält jeden der Bezirke $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ ganz in sich und ist selbst stets in jedem anderen konvexen Bezirk enthalten, der $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ sämtlich in sich aufnimmt.

In der Tat, aus den Regeln 1. und 2. für die m einzelnen Funktionen $H_1, H_2, \dots, H_m$ gehen die entsprechenden Regeln für die Funktion $H(u, v, w)$ hervor. Bedeuten ferner u_1, v_1, w_1 und u_2, v_2, w_2 zwei beliebige Wertsysteme und ist j dann ein solcher Index aus der Reihe $1, 2, \dots, m$, wobei $H_j(u_1 + u_2, v_1 + v_2, w_1 + w_2)$ möglichst groß ausfällt, so haben wir

$$\begin{aligned} H(u_1 + u_2, v_1 + v_2, w_1 + w_2) &= H_j(u_1 + u_2, v_1 + v_2, w_1 + w_2) \\ &\leq H_j(u_1, v_1, w_1) + H_j(u_2, v_2, w_2) \leq H(u_1, v_1, w_1) + H(u_2, v_2, w_2), \end{aligned}$$

also gilt auch die Regel 3. für $H(u, v, w)$ und diese Funktion ist die Stützebenenfunktion eines bestimmten konvexen Bezirks $\mathfrak{R}$. Da nun für $i = 1, 2, \dots, m$ stets $H(u, v, w) \geq H_i(u, v, w)$ ist, so enthält dieser Bezirk $\mathfrak{R}$ jeden der Bezirke $\mathfrak{R}_i$. Ist andererseits $\mathfrak{R}^*$ ein beliebiger konvexer Bezirk, der alle m Bezirke $\mathfrak{R}_i$ enthält, und H^* die Stützebenenfunktion von $\mathfrak{R}^*$, so muß bei jedem System u, v, w stets $H^*(u, v, w) \geq H_i(u, v, w)$ für alle Indizes $i = 1, 2, \dots, m$ und also $H^*(u, v, w) \geq H(u, v, w)$ gelten. Danach enthält dann $\mathfrak{R}^*$ jedesmal den Bezirk $\mathfrak{R} = (\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m)$.

Ein duales Analogon zu der eben dargelegten Tatsache haben wir in folgendem Satze:

Sind $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ eine Reihe von konvexen Körpern, von denen ein jeder den Nullpunkt im Inneren enthält, mit den Distanzfunktionen $f_1, f_2, \dots, f_m$ und bedeutet $f(x, y, z)$ für ein beliebiges System x, y, z jedesmal das Maximum unter den m Werten

$$f_1(x, y, z), f_2(x, y, z), \dots, f_m(x, y, z),$$

so ist $f(x, y, z)$ wieder die Distanzfunktion eines bestimmten konvexen Körpers mit dem Nullpunkte im Inneren, und dieser neue konvexe Körper besteht genau aus allen denjenigen Punkten, die allen m Körpern $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ gemeinsam sind.

§ 12. Die extremen Punkte eines konvexen Bezirks.

Es sei $\mathfrak{R}$ mit der Stützebenenfunktion H ein beliebiger konvexer Bezirk (also ein konvexer Körper oder ein Oval oder eine Strecke oder ein Punkt). Wir bezeichnen einen Punkt p in $\mathfrak{R}$ als einen *extremen Punkt* von $\mathfrak{R}$, wenn p auf keine Weise als inwendiger Punkt einer ganz in $\mathfrak{R}$ enthaltenen Strecke erscheint, d. h. also wenn man nicht in $\mathfrak{R}$ zwei verschiedene Punkte p_1, p_2 und dazu einen Wert $t > 0$ und < 1 finden kann, so daß $p = (1-t)p_1 + tp_2$ gilt.

Wir erkennen sofort: Ein innerer Punkt eines konvexen Körpers, ein inwendiger Punkt eines Ovals in der Ebene des Ovals, weiter ein Punkt einer Strecke, der von ihren Endpunkten verschieden ist, bildet niemals einen extremen Punkt des betreffenden konvexen Bezirks.

Denken wir uns nun p als einen extremen Punkt von $\mathfrak{R}$. Besitzt $\mathfrak{R}$ innere Punkte, d. h. ist $\mathfrak{R}$ ein konvexer Körper, so ist p jedenfalls kein innerer Punkt von $\mathfrak{R}$. Von welcher Art also auch der konvexe Bezirk $\mathfrak{R}$ sein mag, so gehört p jedenfalls der Begrenzung von $\mathfrak{R}$ an und können wir daher wenigstens eine Stützebene durch p an $\mathfrak{R}$ konstruieren. Wir bezeichnen diese Stützebene mit $\{\mathfrak{F}\}$ und es sei (α, β, γ) ihre (äußere) Normale.

Sodann bedeute $\mathfrak{F}$ die gesamte Menge von Punkten, welche $\mathfrak{R}$ in dieser Ebene $\{\mathfrak{F}\}$ liegen hat und wozu insbesondere der Punkt p gehört. Die Eigenschaft eines konvexen Bezirks überträgt sich von $\mathfrak{R}$ auf diesen ebenen Bezirk $\mathfrak{F}$, und wird $\mathfrak{F}$ entweder ein Oval oder eine Strecke oder einen einzelnen Punkt in der Ebene $\{\mathfrak{F}\}$ bedeuten.

Ist $\mathfrak{F}$ ein Oval, so kann p kein inwendiger Punkt dieses Ovals sein, sondern muß notwendig seiner Berandung angehören. Von welcher Art nun auch der Bereich $\mathfrak{F}$ ist, so können wir daher in jedem Falle in der Ebene $\{\mathfrak{F}\}$ wenigstens eine Stützgerade durch p an $\mathfrak{F}$ konstruieren. Wir bezeichnen diese Stützgerade mit $\{\mathfrak{L}\}$ und es sei $(\alpha', \beta', \gamma')$ ihre (äußere) Normale in der Ebene $\{\mathfrak{F}\}$.

Weiter bedeute $\mathfrak{L}$ die Menge der Punkte, welche $\mathfrak{F}$ in dieser Geraden $\{\mathfrak{L}\}$ liegen hat und wozu insbesondere der Punkt p gehört. Die Eigenschaft eines konvexen Bezirks überträgt sich von $\mathfrak{F}$ weiter auf $\mathfrak{L}$ und wird daher $\mathfrak{L}$ entweder eine Strecke oder ein Punkt sein. Im ersteren Falle muß p ein Endpunkt jener Strecke sein, im letzteren $\mathfrak{L}$ sich auf p reduzieren. Es sei endlich $(\alpha'', \beta'', \gamma'')$ von den beiden Richtungen in der Geraden $\{\mathfrak{L}\}$ im ersteren Falle diejenige, welche vom Punkte p aus von der Strecke $\{\mathfrak{L}\}$ fortführt, im zweiten Falle eine beliebige dieser zwei Richtungen. Die drei Richtungen (α, β, γ) , $(\alpha', \beta', \gamma')$, $(\alpha'', \beta'', \gamma'')$ stehen aufeinander senkrecht, so daß zu jedem extremen Punkte von $\mathfrak{R}$ in der hier dargelegten Weise wenigstens eine orthogonale Substitution gehört.

Andererseits leuchtet ein, daß zu einer jeden beliebigen linearen Substitution, insbesondere zu einer orthogonalen Substitution

$$(S) \quad \xi = \alpha''x + \beta''y + \gamma''z, \quad \eta = \alpha'x + \beta'y + \gamma'z, \quad \zeta = \alpha x + \beta y + \gamma z$$

stets ein bestimmter Punkt in $\mathfrak{R}$ gehört, der unter allen Punkten von $\mathfrak{R}$ zunächst einen möglichst großen Wert von ξ , nächst dem einen möglichst großen Wert von η und nächst dem einen möglichst großen Wert von ζ besitzt. Dieser Punkt p ist dann jedesmal ein extremer Punkt von $\mathfrak{R}$. Denn gehörte p einer Strecke $p_1 p_2$ an, wobei p_1 und p_2 zwei von p verschiedene Punkte aus $\mathfrak{R}$ wären, so kann zunächst ξ weder für p_1 , noch für p_2 größer wie für p sein und müßte daher für beide Punkte denselben Wert wie für p haben, weiter müßte η und endlich auch ζ für p_1 und p_2 denselben Wert wie für p haben, was unmöglich ist. Wir bezeichnen den in Frage kommenden Punkt p als den zur Substitution (S) gehörenden extremen Punkt von $\mathfrak{R}$. Sehen wir nun zu, wie dieser zu (S) gehörende extreme Punkt von $\mathfrak{R}$ sich mittels der Stützebenenfunktion $H(u, v, w)$ des Bezirks $\mathfrak{R}$ charakterisieren läßt.

Indem wir uns von vornherein die orthogonale Koordinatentransformation (S) vorgenommen denken, können wir der Einfachheit wegen annehmen, es handle sich speziell um den extremen Punkt, der zu den ursprünglichen Koordinaten, d. h. der Substitution

$$(46) \quad \xi = x, \quad \eta = y, \quad \zeta = z$$

gehört. Alsdann ist zunächst für (α, β, γ) die Richtung der positiven z -Achse einzuführen und also unter $\{\mathfrak{F}\}$ die Stützebene $z = H(0, 0, 1)$ an $\mathfrak{R}$ zu verstehen. Ermitteln wir nun den Bereich $\mathfrak{F}$ der Punkte von $\mathfrak{R}$ in dieser Ebene.

Nach den Ausführungen in § 8 und § 11 werden in dieser Ebene zu $\mathfrak{R}$ alle diejenigen Punkte x, y, z gehören, für deren Koordinaten x, y bei beliebigen Werten u, v, w stets

$$ux + vy \leq H(u, v, w) - wH(0, 0, 1)$$

gilt. Wir setzen zur Abkürzung

$$(47) \quad H(u, v, w) - wH(0, 0, 1) = H'(u, v, w);$$

diese Funktion H' genügt dann ebenso wie H den Bedingungen 1., 2., 3. in § 11. Dabei wird $H'(0, 0, 1) = 0$, $H'(0, 0, -1) \geq 0$, also allgemein

$$H'(0, 0, t) = 0, \quad H'(0, 0, -t) \geq 0,$$

wenn $t > 0$ ist. Aus

$$H'(u, v, w_0) + H'(0, 0, w - w_0) \geq H'(u, v, w)$$

geht nunmehr, wenn $w > w_0$ ist, stets $H'(u, v, w_0) \geq H'(u, v, w)$ hervor. Danach nimmt die Funktion $H'(u, v, w)$ mit wachsendem w bei festen u, v niemals zu. Andererseits ist stets

$$H'(u, v, w) + H'(-u, -v, 0) \geq H'(0, 0, w) \geq 0,$$

$$H'(u, v, w) \geq -H'(-u, -v, 0),$$

so daß $H'(u, v, w)$ bei festen Werten u, v nicht unter eine bestimmte Grenze sinken kann. Mithin wird die Funktion $H'(u, v, w)$ bei festen Werten u, v für ein unbegrenzt wachsendes w , ohne jemals zuzunehmen, nach einer endlichen unteren Grenze konvergieren (die sie auch schon von einem endlichen Werte w an erreichen kann); diese Grenze $H'(u, v, +\infty)$ wollen wir schlechthin mit $H'(u, v)$ bezeichnen. Aus den Eigenschaften der Funktion $H'(u, v, w)$ gehen folgende Umstände für die Funktion $H'(u, v)$ hervor:

$$1. \quad H'(u, v) + H'(-u, -v) \geq 0, \quad H'(0, 0) = 0;$$

$$2. \quad H'(tu, tv) = tH'(u, v), \quad \text{wenn } t > 0 \text{ ist};$$

$$3. \quad H'(u_1, v_1) + H'(u_2, v_2) \geq H'(u_1 + u_2, v_1 + v_2).$$

Die Punkte x, y, z des Bezirks $\mathfrak{F}$ sind nun völlig charakterisiert durch $z = H(0, 0, 1)$ und die Ungleichungen

$$(48) \quad ux + vy \leq H'(u, v)$$

für alle möglichen Wertsysteme u, v .

Weiter ist für $(\alpha', \beta', \gamma')$ die Richtung der positiven y -Achse einzuführen und unter $\{\mathfrak{Q}\}$ die Stützgerade $y = H'(0, 1)$ an $\mathfrak{F}$ in der Ebene $\{\mathfrak{F}\}$ zu verstehen. Durch analoge Betrachtungen, wie sie soeben angestellt wurden, erkennen wir, daß die Funktion

$$(49) \quad H''(u, v) - vH'(0, 1) = H''(u, v)$$

bei festem Werte u für ein unbegrenzt wachsendes v , ohne jemals zuzunehmen, nach einer endlichen unteren Grenze $H''(u, +\infty)$ konvergiert, die wir einfach mit $H''(u)$ bezeichnen. Alsdann gilt

$$H''(0) = 0; \quad H''(tu) = tH''(u), \quad \text{wenn } t > 0 \text{ ist}; \quad H''(1) + H''(-1) \geq 0;$$

und finden wir den Bereich $\mathfrak{L}$ der Punkte von $\mathfrak{F}$ in der Geraden $\{\mathfrak{L}\}$ durch

$$(50) \quad z = H(0, 0, 1), \quad y = H'(0, 1), \quad -H''(-1) \leq x \leq H''(1)$$

bestimmt.

Endlich soll noch $(\alpha'', \beta'', \gamma'')$ die Richtung der positiven x -Achse bedeuten und ist der zur Substitution (46) gehörende extreme Punkt p von $\mathfrak{R}$ völlig charakterisiert durch die Koordinatenwerte

$$(51) \quad z = H(0, 0, 1), \quad y = H'(0, 1), \quad x = H''(1).$$

Zu jedem konvexen Bezirk $\mathfrak{R}$ gehören in der hier dargelegten Weise gewisse Bezirke $\mathfrak{F}$ und $\mathfrak{L}$ in Ebenen und Geraden und gewisse extreme Punkte p . Der Bezirk $\mathfrak{R}$ besteht alsdann 1) aus seinen inneren Punkten, falls $\mathfrak{R}$ ein konvexer Körper ist, 2) aus den inwendigen Punkten der Bezirke $\mathfrak{F}$, soweit diese Ovale sind, 3) aus den inwendigen Punkten der Bezirke $\mathfrak{L}$, soweit diese Strecken sind, endlich 4) aus den extremen Punkten p .

Ist $\mathfrak{R}$ ein konvexer Körper, p ein beliebiger extremer Punkt von $\mathfrak{R}$, so finden wir die inneren Punkte von $\mathfrak{R}$ auf den Strecken von p nach den anderen Punkten der Begrenzung von $\mathfrak{R}$ gelegen. Ist $\mathfrak{F}$ ein Oval, p ein beliebiger extremer Punkt von $\mathfrak{F}$, so liegen die inwendigen Punkte von $\mathfrak{F}$ auf den Strecken von p nach den anderen Punkten der Berandung von $\mathfrak{F}$. Ist $\mathfrak{L}$ eine Strecke, so sind die Endpunkte ihre extremen Punkte.

Überblicken wir diese Tatsachen, so erkennen wir, daß in einem konvexen Bezirk $\mathfrak{R}$ ein beliebiger Punkt stets entweder selbst ein extremer Punkt von $\mathfrak{R}$ ist oder einer Strecke oder einem Dreieck oder einem Tetraeder angehört, deren zwei, drei, vier Eckpunkte lauter extreme Punkte von $\mathfrak{R}$ sind.

Ein konvexer Bezirk $\mathfrak{R}$ ist auf diese Weise durch die Gesamtheit seiner extremen Punkte völlig bestimmt als der kleinste konvexe Bezirk, der alle diese Punkte aufnimmt. *Ein konvexer Bezirk mit einer endlichen Anzahl von extremen Punkten* stellt danach gemäß den Definitionen in § 4 stets entweder ein Polyeder oder ein Polygon oder eine Strecke oder einen Punkt vor. Umgekehrt sind bei den eben genannten Bereichen nach einer Bemerkung in § 4 allein die Eckpunkte extreme Punkte und also die extremen Punkte in der Tat jedesmal nur in endlicher Anzahl vorhanden.

Wir beweisen noch den folgenden Satz:

Ist $\mathfrak{R}$ ein konvexer Körper mit dem Nullpunkte o im Inneren und ε eine beliebige positive Größe, so können wir stets ein Polyeder $\mathfrak{Q}$ bestimmen, dessen Ecken sämtlich extreme Punkte von $\mathfrak{R}$ sind, welches daher ganz in $\mathfrak{R}$ enthalten ist, und so daß andererseits das Polyeder $(1 + \varepsilon)\mathfrak{Q}$ den Körper $\mathfrak{R}$ ganz enthält.

In der Tat, wir bestimmen zunächst nach § 5 ein Polyeder Ω^* irgendwie so, daß Ω^* in $\mathfrak{K}$ und $\mathfrak{K}$ in $(1+\varepsilon)\Omega^*$ enthalten ist. Jeden einzelnen Eckpunkt q^* von Ω^* können wir, wenn er nicht selbst ein extremer Punkt ist, als Punkt einer Strecke oder eines Dreiecks oder eines Tetraeders darstellen, deren zwei, drei, vier Eckpunkte extreme Punkte von $\mathfrak{K}$ sind. Es bedeute sodann Ω den kleinsten konvexen Bezirk, welcher sämtliche verschiedenen hierbei herangezogenen extremen Punkte von $\mathfrak{K}$ in sich aufnimmt, so stellt dieser Bezirk Ω , der natürlich nicht ganz in eine Ebene fällt, ein Polyeder mit diesen extremen Punkten von $\mathfrak{K}$ als Eckpunkten vor. Dabei enthält nun Ω jeden Eckpunkt von Ω^* und also das ganze Polyeder Ω^* und wird daher $(1+\varepsilon)\Omega$ das Polyeder $(1+\varepsilon)\Omega^*$ und somit den Körper $\mathfrak{K}$ enthalten. Das Polyeder Ω entspricht danach allen in dem obigen Satze gestellten Anforderungen.

§ 13. Projektionsraum eines konvexen Körpers von einem Punkte seiner Begrenzung. Extreme Stützebenen.

Wir haben jetzt einige Bemerkungen über das Verhalten eines konvexen Körpers in der Umgebung eines Punktes seiner Begrenzung anzuschließen. Es sei $\mathfrak{K}$ ein konvexer Körper mit der Stützebenenfunktion $H(u, v, w)$ und p mit den Koordinaten x_0, y_0, z_0 ein beliebiger Punkt seiner Begrenzung. Wir denken uns jeden Strahl von p aus konstruiert, der durch innere Punkte von $\mathfrak{K}$ geht. Die gesamte Menge der Punkte auf allen diesen Strahlen, mit Ausschluß des Punktes p selbst, werde mit $\mathfrak{K}'$ bezeichnet. Da um jeden inneren Punkt von $\mathfrak{K}$ eine Kugel, aus lauter inneren Punkten von $\mathfrak{K}$ bestehend, konstruiert werden kann, besitzt die Menge $\mathfrak{K}'$ offenbar nur innere Punkte und enthält also keinen Punkt ihrer Begrenzung. Es sei sodann $\mathfrak{K}$ diejenige abgeschlossene Menge, welche aus $\mathfrak{K}'$ und der ganzen Begrenzung von $\mathfrak{K}'$ besteht; diese Menge $\mathfrak{K}$ nennen wir den *Projektionsraum* des Körpers $\mathfrak{K}$ von p aus. Wir denken uns ferner um p als Mittelpunkt eine Kugel vom Radius 1 konstruiert, und das Gebiet P , welches $\mathfrak{K}$ mit dieser Kugel gemein hat, heiße der *Projektionssektor* des Körpers $\mathfrak{K}$ von p aus.

Sind r_1' und r_2' zwei Punkte aus $\mathfrak{K}'$ auf zwei verschiedenen Strahlen von p aus, so können wir auf diesen Strahlen irgend zwei innere Punkte q_1 und q_2 von $\mathfrak{K}$ angeben. Jeder Punkt der Strecke q_1q_2 ist dann ebenfalls ein innerer Punkt von $\mathfrak{K}$ und gehört dadurch jeder Punkt der Strecke $r_1'r_2'$ wieder zu $\mathfrak{K}'$. Die Menge $\mathfrak{K}'$ besitzt also die Eigenschaft 1a) eines konvexen Bezirks.

Zu jedem Punkte r von $\mathfrak{K}$ gibt es, wenn er nicht zu $\mathfrak{K}'$ gehört, in beliebig kleiner Entfernung von ihm Punkte r' aus $\mathfrak{K}'$. Sind r_1, r_2 zwei

Punkte aus $\mathfrak{R}$, so wird man daher zwei Punkte r_1', r_2' aus $\mathfrak{R}'$ angeben können, so daß die Entfernungen $r_1 r_1', r_2 r_2'$ unterhalb einer beliebig kleinen Größe liegen, und da die Strecke $r_1' r_2'$ ganz zu $\mathfrak{R}'$ gehört, existiert also auch zu jedem Punkte der Strecke $r_1 r_2$ in beliebiger Umgebung ein Punkt von $\mathfrak{R}'$, ist also jeder Punkt dieser Strecke ein Punkt von $\mathfrak{R}'$ oder doch eine Häufungsstelle von $\mathfrak{R}'$, mithin stets in $\mathfrak{R}$ enthalten. Es besitzt also auch die Menge $\mathfrak{R}$ die Eigenschaft 1a) eines konvexen Bezirks; dabei besteht $\mathfrak{R}$ aus lauter Strahlen von p aus. Der Raum $\mathfrak{R}$ besitzt weiter die Eigenschaft 1c) eines konvexen Bezirks, abgeschlossen zu sein, und enthält den ganzen Körper $\mathfrak{R}$, da jeder Punkt der Begrenzung von $\mathfrak{R}$ eine Häufungsstelle von inneren Punkten von $\mathfrak{R}$ ist.

Die Menge $\mathfrak{R}'$ besteht nun genau aus allen inneren Punkten von $\mathfrak{R}$. Denn ist r irgendein innerer Punkt von $\mathfrak{R}$, so können wir vier nicht in einer Ebene gelegene Punkte r_1, r_2, r_3, r_4 aus $\mathfrak{R}$ so wählen, daß r im Inneren des Tetraeders (r_1, r_2, r_3, r_4) liegt. Weiter können wir in beliebiger Nähe von jedem dieser Punkte einen Punkt aus $\mathfrak{R}'$ finden und daher auch vier Punkte r_1', r_2', r_3', r_4' aus $\mathfrak{R}'$ wählen, so daß r dem Bezirke (r_1', r_2', r_3', r_4') angehört; alsdann aber gehört nach dem vorhin Auseinandergesetzten r selbst zu $\mathfrak{R}'$. Andererseits kann $\mathfrak{R}'$, wie wir bereits gesehen haben, keinen Punkt der Begrenzung von $\mathfrak{R}$ enthalten.

Die Eigenschaften 1a), 1c), 2) eines konvexen Körpers übertragen sich von dem Projektionsraume $\mathfrak{R}$ sofort auf den Projektionssektor P ; dieser besitzt schließlich noch die Eigenschaft 1b) eines konvexen Bezirks, ganz im Endlichen zu liegen, und stellt somit in der Tat einen konvexen Körper vor. Nach § 8 ist nun ein konvexer Körper völlig bestimmt durch die Bedingungen seiner Stützebenen, und zwar geht aus den Ausführungen bei (21) dort hervor, daß zur Charakterisierung des Körpers nicht durchaus seine sämtlichen Stützebenen erforderlich sind, sondern daß dazu bereits die Bedingungen jeder solchen Menge unter den Stützebenen völlig hinreichen, welche für sich schon die Eigenschaft hat, jeden einzelnen Punkt der Begrenzung von $\mathfrak{R}$ aufzunehmen. Nun baut sich der Projektionssektor P aus lauter Radien der Kugel

$$(52) \quad (x - x_0)^2 + (y - y_0)^2 + (z - z_0)^2 \leq 1$$

auf. Durch jeden Endpunkt eines solchen Radius geht eine Tangentialebene dieser Kugel, welche zugleich eine Stützebene an P bildet. Ist aber r ein Zwischenpunkt eines solchen Radius und gehört der Begrenzung von P an, so kann eine Stützebene durch r an P nicht Punkte dieses Radius trennen und muß also durch p gehen, ist somit eine Stützebene durch p an P und dann weiter auch an $\mathfrak{R}$ und endlich an $\mathfrak{R}$ selbst. Andererseits hat jede Stützebene durch p an $\mathfrak{R}$ den Raum $\mathfrak{R}'$ ganz auf

einer Seite liegen und ist daher zugleich auch Stützebene durch p an $\mathfrak{R}$ und an P .

Hiernach wird nun der Projektionssektor P bereits völlig charakterisiert sein durch die Ungleichung (52), welche die Bedingungen aller Stützebenen an diese Kugel zusammenfaßt, und zudem durch die Bedingungen aller Stützebenen

$$(53) \quad \alpha x + \beta y + \gamma z \leq H(\alpha, \beta, \gamma)$$

an $\mathfrak{R}$, welche durch p gehen, also die Beziehung

$$(54) \quad \alpha x_0 + \beta y_0 + \gamma z_0 = H(\alpha, \beta, \gamma)$$

erfüllen. Nach der Art, wie der Projektionsraum $\mathfrak{R}$ und der Projektionssektor P sich gegenseitig bestimmen, wird dann der Projektionsraum $\mathfrak{R}$ schon allein durch die letzteren Bedingungen (53), (54), also durch die Bedingungen der sämtlichen vorhandenen Stützebenen durch p an $\mathfrak{R}$ charakterisiert sein. Die Menge $\mathfrak{R}'$ endlich bildet, wie wir sahen, das ganze Innere des Raumes $\mathfrak{R}$. —

Wir nennen zwei Richtungen unabhängig, wenn sie weder identisch noch einander entgegengesetzt sind, drei Richtungen unabhängig, wenn sie nicht sämtlich einer einzigen Ebene angehören. Wir bezeichnen einen Punkt p der Begrenzung von $\mathfrak{R}$ als einen *Flächenpunkt* von $\mathfrak{R}$, wenn durch p nur eine einzige Stützebene an $\mathfrak{R}$ geht, als einen *Kantenpunkt* von $\mathfrak{R}$, wenn durch p mehrere Stützebenen an $\mathfrak{R}$ gehen, aber deren Normalenrichtungen sämtlich einer Ebene angehören, als einen *Eckpunkt* von $\mathfrak{R}$, wenn durch p sich drei solche Stützebenen an $\mathfrak{R}$ legen lassen, deren Normalen nicht sämtlich einer einzigen Ebene angehören.

Gilt in $\mathfrak{R}$ durchweg $\varphi \leq 0$, wobei $\varphi = 0$ die Gleichung einer Ebene, aber nicht durchaus einer Stützebene an $\mathfrak{R}$ vorstellt, so wollen wir den durch $\varphi \leq 0$ bestimmten Bereich als einen *Halbraum* um $\mathfrak{R}$ bezeichnen. Sind in diesem Sinne $\varphi_1 \leq 0$, $\varphi_2 \leq 0$ zwei Halbräume um $\mathfrak{R}$ und sind c_1 , c_2 zwei positive Konstante, so gilt in $\mathfrak{R}$ auch stets $\varphi = c_1 \varphi_1 + c_2 \varphi_2 \leq 0$ und stellt also diese Bedingung $\varphi \leq 0$ jedesmal wieder einen Halbraum um $\mathfrak{R}$ dar. Wir bezeichnen nun einen Halbraum $\varphi \leq 0$ um $\mathfrak{R}$ als einen *extremen Halbraum* um $\mathfrak{R}$, wenn es nicht möglich ist, zwei *verschiedene* Halbräume $\varphi_1 \leq 0$, $\varphi_2 \leq 0$ um $\mathfrak{R}$ und zwei positive Konstanten c_1 , c_2 so zu wählen, daß $\varphi = c_1 \varphi_1 + c_2 \varphi_2$ ist.

Wir denken uns jetzt der Einfachheit wegen den Anfangspunkt o der Koordinaten im Inneren von $\mathfrak{R}$ gelegen, um gemäß § 8 den polaren Körper $\mathfrak{H}$ von $\mathfrak{R}$ einzuführen. Es gilt dann stets $H(u, v, w) > 0$, wenn $u, v, w \neq 0, 0, 0$ sind. Die Bedingung eines jeden des Nullpunkt im Inneren einschließenden Halbraumes können wir in die Form

$$(55) \quad \psi = ux + xy + wz - 1 \leq 0$$

setzen, und ein solcher „Halbraum (u, v, w) “ enthält den ganzen Körper $\mathfrak{R}$, jedesmal dann, wenn $H(u, v, w) \leq 1$ ist. Andererseits besteht der polare Körper $\mathfrak{H}$ aus allen Punkten mit Koordinaten u, v, w , wobei $H(u, v, w) \leq 1$ ist, und entsprechen jetzt die extremen Halbräume (u, v, w) um $\mathfrak{R}$ offenbar genau den extremen Punkten (u, v, w) des Körpers $\mathfrak{H}$. Da letztere Punkte notwendig der Begrenzung von $\mathfrak{H}$ angehören, also für sie stets $H(u, v, w) = 1$ ist, sehen wir zunächst, daß die Bedingung eines extremen Halbraumes um $\mathfrak{R}$ jedenfalls zugleich die Bedingung einer *Stützebene* an $\mathfrak{R}$ vorstellt; die bezüglichen Stützebenen bezeichnen wir als *extreme Stützebenen* an $\mathfrak{R}$.

Dem Punkte $p(x_0, y_0, z_0)$ auf der Begrenzung von $\mathfrak{R}$ entspricht dualistisch die Stützebene

$$(56) \quad x_0 u + y_0 v + z_0 w - 1 \leq 0$$

an $\mathfrak{H}$. Die Menge der Punkte von $\mathfrak{H}$ in dieser Ebene wollen wir mit $\mathfrak{D}(p)$ bezeichnen; jeder Stützebene durch p an $\mathfrak{R}$ mit einer Bedingung $u_0 x + v_0 y + w_0 z - 1 \leq 0$ entspricht in $\mathfrak{D}(p)$ ein Punkt (u_0, v_0, w_0) . Je nachdem nun p ein Flächen-, Kanten-, Eckpunkt von $\mathfrak{R}$ ist, wird $\mathfrak{D}(p)$ einen einzelnen Punkt oder eine Strecke oder ein Oval vorstellen. Auf Grund der oben abgeleiteten Resultate erkennen wir nun die folgenden Umstände.

Ist zunächst p ein *Flächenpunkt* von $\mathfrak{R}$ und $\psi \leq 0$ die Bedingung der einzigen alsdann vorhandenen Stützebene durch p an $\mathfrak{R}$, so wird der Projektionsraum $\mathfrak{R}$ des Körpers $\mathfrak{R}$ von p aus der ganze *Halbraum* $\psi \leq 0$.

Ist zweitens p ein *Kantenpunkt* von $\mathfrak{R}$, so gibt es zwei extreme Stützebenen durch p an $\mathfrak{R}$, deren Bedingungen $\psi_1 \leq 0$, $\psi_2 \leq 0$ seien, und diese haben dann die Bedingung jeder anderen durch p an $\mathfrak{R}$ vorhandenen Stützebene in einer Form $(1-t)\psi_1 + t\psi_2 \leq 0$, $0 < t < 1$ zur Folge. Der Projektionsraum $\mathfrak{R}$ des Körpers $\mathfrak{R}$ von p aus wird hier der durch die beiden Ungleichungen $\psi_1 \leq 0$, $\psi_2 \leq 0$ dargestellte *Keil*.

Wir schalten jetzt eine Bemerkung in bezug auf die ebenen Figuren ein. Wir können die hier für den Raum und konvexe Körper eingeführten Begriffe sinngemäß auf die Ebene und Ovale übertragen. Ist $\mathfrak{S}$ ein ebenes Oval, f ein Punkt seiner Berandung, so nennen wir f einen *Linienpunkt* oder einen *Eckpunkt* von $\mathfrak{S}$, je nachdem in der Ebene von $\mathfrak{S}$ durch f nur eine oder mehrere Stützgeraden an $\mathfrak{S}$ gehen. Wir bilden sodann in der gehörigen Weise die Begriffe der extremen Halbebene um $\mathfrak{S}$ und der extremen Stützgeraden an $\mathfrak{S}$ und erkennen: Im Falle eines Linienpunktes f ist die einzige Stützgerade durch f an $\mathfrak{S}$ eine *extreme Stützgerade* an $\mathfrak{S}$, im Falle eines Eckpunktes f gibt es zwei verschiedene *extreme Stützgeraden* durch f an $\mathfrak{S}$, aus welchen dann die Bedingungen aller übrigen Stützgeraden der „Eckstützgeraden“ durch f an $\mathfrak{S}$ folgen.

Jetzt nehmen wir drittens p als einen *Eckpunkt* von $\mathfrak{R}$ an, wobei dann $\mathfrak{D}(p)$ ein Oval vorstellt. Diejenigen Stützebenen durch p an $\mathfrak{R}$, welche den *inwendigen* Punkten dieses Ovals entsprechen, bezeichnen wir als *Eckstützebenen* durch p an $\mathfrak{R}$. Zu jedem inwendigen Punkte b von $\mathfrak{D}(p)$ können wir drei nicht in gerader Linie gelegene Punkte b_1, b_2, b_3 in $\mathfrak{D}(p)$ derart annehmen, daß b im Inneren des Dreiecks (b_1, b_2, b_3) liegt. Die charakteristische Eigenschaft einer Eckstützebene durch p an $\mathfrak{R}$ ist danach die, daß man ihre Bedingung in eine Form

$$(57) \quad \varphi = c_1 \varphi_1 + c_2 \varphi_2 + c_3 \varphi_3 \leq 0$$

setzen kann, so daß $\varphi_1 \leq 0, \varphi_2 \leq 0, \varphi_3 \leq 0$ die Bedingungen dreier Stützebenen durch p an $\mathfrak{R}$ mit *unabhängigen* Normalenrichtungen und c_1, c_2, c_3 positive Konstanten sind.

Es bedeute in dieser Weise (57) eine Eckstützebene durch p an $\mathfrak{R}$. Alsdann sind $\varphi_1, \varphi_2, \varphi_3$ drei unabhängige homogene lineare Ausdrücke in $x - x_0, y - y_0, z - z_0$, so daß umgekehrt letztere Differenzen durch $\varphi_1, \varphi_2, \varphi_3$ ausgedrückt werden können. Im Projektionsraume $\mathfrak{R}$ des Körpers $\mathfrak{R}$ von p aus gilt jedenfalls

$$(58) \quad \varphi_1 \leq 0, \varphi_2 \leq 0, \varphi_3 \leq 0.$$

Fordern wir gleichzeitig $\varphi = 0$, so muß wegen (57) notwendig $\varphi_1 = 0, \varphi_2 = 0, \varphi_3 = 0$ gelten. Die Ebene $\varphi = 0$ enthält also vom Raume $\mathfrak{R}$ allein den Punkt p , während im übrigen in $\mathfrak{R}$ stets $\varphi < 0$ ist. Bezeichnen wir den Schnitt von $\mathfrak{R}$ mit der Ebene $\varphi = -1$ mit $\mathfrak{S}$, so gilt wegen (57) und (58) in $\mathfrak{S}$:

$$(59) \quad -\frac{1}{c_1} \leq \varphi_1 \leq 0, \quad -\frac{1}{c_2} \leq \varphi_2 \leq 0, \quad -\frac{1}{c_3} \leq \varphi_3 \leq 0,$$

und danach bestehen auch für die Beträge von $x - x_0, y - y_0, z - z_0$ bei den Punkten in $\mathfrak{S}$ obere Grenzen. Der Bereich $\mathfrak{S}$ ist also ganz im Endlichen gelegen; der Raum $\mathfrak{R}$ ist dann der Kegel aller Strahlen von p durch die Punkte von $\mathfrak{S}$, und enthält gewiß drei nicht in gerader Linie gelegene Punkte. Von $\mathfrak{R}$ übernimmt $\mathfrak{S}$ ferner die Eigenschaften 1a) und 1c) eines konvexen Bezirks und stellt nun $\mathfrak{S}$ ein *Oval* in der Ebene $\varphi = -1$ vor.

Als Stützebenen durch p an $\mathfrak{R}$ oder an $\mathfrak{R}$ haben wir dann außer der Ebene $\varphi = 0$ jede Ebene durch p , welche die Ebene $\varphi = -1$ in einer solchen Geraden schneidet, die das Oval $\mathfrak{S}$ überhaupt nicht trifft oder eine Stützgerade an $\mathfrak{S}$ ist. Aus den entsprechenden Tatsachen über die Stützgeraden eines Ovals entnehmen wir, daß eine Stützebene durch p an $\mathfrak{R}$ dann und nur dann eine *extreme* Stützebene an $\mathfrak{R}$ ist, wenn sie durch eine extreme Stützgerade an $\mathfrak{S}$ in $\varphi = -1$ führt. Ferner erkennen wir als die Eckstützebenen durch p an $\mathfrak{R}$ diejenigen Ebenen durch p , welche $\mathfrak{S}$ überhaupt nicht treffen. Ist q ein von p verschiedener Punkt, der

weder in $\mathfrak{R}$ noch in dem zu $\mathfrak{R}$ in bezug auf p symmetrischen Kegel liegt, so gehen durch die Gerade pq offenbar noch unendlich viele dieser Eckstützebenen.

Der Projektionssektor P wird nun im Falle eines Flächenpunktes eine Halbkugel, im Falle eines Kantenpunktes ein Keilsektor, im Falle eines Eckpunktes ein Kegelsektor.

Nach dem am Schlusse von § 12 bewiesenen Satze können wir, wenn ε eine beliebige positive GröÙe bedeutet, stets zum Körper $\mathfrak{H}$ ein Polyeder $\mathfrak{P}$ bestimmen, dessen Ecken lauter extreme Punkte von $\mathfrak{H}$ sind und so daß $(1 + \varepsilon)\mathfrak{P}$ den Körper $\mathfrak{H}$ ganz enthält. Dabei ist der Nullpunkt o ein innerer Punkt von $(1 + \varepsilon)\mathfrak{P}$ und also auch von $\mathfrak{P}$. Es bedeute $\mathfrak{Q}$ den zu $\mathfrak{P}$ polaren Körper, so ist nunmehr $\mathfrak{Q}$ ein Polyeder, bei welchem die Ebenen der einzelnen Seitenflächen lauter extreme Stützebenen an $\mathfrak{R}$ sind, und enthält $\mathfrak{Q}$ den Körper $\mathfrak{R}$. Andererseits wird der zu $(1 + \varepsilon)\mathfrak{P}$ polare Körper das Polyeder $\frac{1}{1 + \varepsilon}\mathfrak{Q}$ sein und ganz in $\mathfrak{R}$ enthalten sein.

Wir gelangen damit zu folgendem Satze:

Ist $\mathfrak{R}$ ein konvexer Körper mit dem Nullpunkte im Inneren und ε eine beliebige positive GröÙe, so können wir stets ein Polyeder $\mathfrak{Q}$ bestimmen, bei welchem die Ebenen der Seitenflächen lauter extreme Stützebenen an $\mathfrak{R}$ sind und so daß andererseits $\mathfrak{Q}$ ganz in $(1 + \varepsilon)\mathfrak{R}$ enthalten ist.

§ 14. Tangentialebenen.

Es sei $\mathfrak{R}$ ein konvexer Körper und

$$(60) \quad \psi = \alpha x + \beta y + \gamma z - H(\alpha, \beta, \gamma) \leq 0$$

die Bedingung der Stützebene an $\mathfrak{R}$ mit der beliebigen Richtung (α, β, γ) als (äußerer) Normale. So lange d positiv ist und eine gewisse GröÙe $(H(\alpha, \beta, \gamma) + H(-\alpha, -\beta, -\gamma))$ nicht erreicht, trifft die parallele Ebene $\psi = -d$ den Körper $\mathfrak{R}$, ohne an ihn Stützebene zu sein, geht also durch innere Punkte von $\mathfrak{R}$ und hat mit $\mathfrak{R}$ einen ebenen Bereich gemein, der von $\mathfrak{R}$ die Eigenschaft eines konvexen Bezirks übernimmt und gewiß drei nicht in gerader Linie gelegene Punkte aufweist, also ein Oval vorstellt. Es sei $\varrho(d)$ das Maximum unter den Radien der ganz in diesem Oval enthaltenen Kreise. Wir bezeichnen die Stützebene $\psi \leq 0$ an $\mathfrak{R}$ als eine *Tangentialebene* an $\mathfrak{R}$, wenn der Quotient $\frac{\varrho(d)}{d}$ für ein nach Null abnehmendes d über jede Grenze hinauswächst. Wir können nun die folgende Tatsache beweisen:

Die extremen Stützebenen an $\mathfrak{R}$ und allein diese Stützebenen sind Tangentialebenen an $\mathfrak{R}$.

Fassen wir einen beliebigen Punkt p der Begrenzung von $\mathfrak{R}$ ins Auge.

Ist zunächst p ein *Flächenpunkt* von $\mathfrak{K}$, so ist die alsdann vorhandene einzige Stützebene durch p an $\mathfrak{K}$ nach § 13 eine extreme Stützebene an $\mathfrak{K}$, und wie wir jetzt zeigen wollen, zugleich Tangentialebene an $\mathfrak{K}$ in dem eben festgelegten Sinne. Bedeutet (60) die Bedingung dieser Stützebene, so ist der Projektionsraum $\mathfrak{H}$ des Körpers $\mathfrak{K}$ von p aus hier der ganze Halbraum $\psi \leq 0$. Die Parallelebene $\psi = -1$ fällt daher vollständig ins *Innere* dieses Raumes $\mathfrak{H}$ und können wir drei innere Punkte r_1, r_2, r_3 von $\mathfrak{H}$ in dieser Ebene $\psi = -1$ so annehmen, daß das Dreieck (r_1, r_2, r_3) einen Kreis von beliebig großem Radius ϱ_0 enthält. Als dann können wir auf den Strecken pr_1, pr_2, pr_3 irgendwelche inneren Punkte q_1, q_2, q_3 von $\mathfrak{K}$ finden; ist für diese Punkte $\psi = -d_1, -d_2, -d_3$, so gehören bei jedem positiven Werte d , der $\leq d_1, \leq d_2, \leq d_3$ ist, alle drei Punkte

$$(1-d)p + dr_1, \quad (1-d)p + dr_2, \quad (1-d)p + dr_3$$

in der Ebene $\psi = -d$ und daher auch das ganze Dreieck mit letzteren Punkten als Eckpunkten dem Inneren von $\mathfrak{K}$ an. Dieses dem Dreieck (r_1, r_2, r_3) in bezug auf p homothetische Dreieck enthält einen Kreis vom Radius $d\varrho_0$ und ist also dabei $\frac{\varrho(d)}{d} \geq \varrho_0$. Danach muß in der Tat der Quotient $\frac{\varrho(d)}{d}$ für ein nach Null abnehmendes d über jede Grenze hinauswachsen.

Ist zweitens p ein *Kantenpunkt* von $\mathfrak{K}$, so ist der Projektionsraum $\mathfrak{H}$ des Körpers $\mathfrak{K}$ von p aus nach § 13 ein von zwei extremen Stützebenen $\psi_1 \leq 0, \psi_2 \leq 0$ an $\mathfrak{K}$ gebildeter Keil. Als dann sind diese beiden extremen Stützebenen Tangentialebenen an $\mathfrak{K}$. Denn betrachten wir eine dieser Ebenen, $\psi_1 = 0$. Der Keil $\mathfrak{H}$ schneidet aus der zu ihr parallelen Ebene $\psi_1 = -1$ eine Halbebene heraus. Wir können drei innere Punkte r_1, r_2, r_3 von $\mathfrak{H}$ in dieser Halbebene so annehmen, daß das Dreieck (r_1, r_2, r_3) einen Kreis von beliebig großem Radius enthält, und die weiteren Schlüsse genau wie vorhin einrichten, um für $\psi_1 = 0$ den Charakter einer Tangentialebene an $\mathfrak{K}$ einzusehen. — Daß die übrigen Stützebenen durch p an $\mathfrak{K}$, welche mit dem Keile $\mathfrak{H}$ nur die Kante $\psi_1 = 0, \psi_2 = 0$ gemein haben, *nicht* Tangentialebenen an $\mathfrak{K}$ vorstellen, leuchtet ebenso einfach ein.

Analoge Bemerkungen wie hier zu den Flächen- und Kantenpunkten eines konvexen Körpers können wir zu den Linien- und Eckpunkten eines ebenen Ovals machen, und wir kommen dadurch zu folgendem Satze: Es sei $\mathfrak{S}$ ein ebenes Oval, f ein Punkt seiner Berandung, ft eine Stützgerade an $\mathfrak{S}$ und bezeichnen wir die Länge der zu ft parallelen Sehne von $\mathfrak{S}$ in einem senkrechten Abstände d' von ft mit $\varrho'(d')$, so wächst für ein nach Null abnehmendes d' der Quotient $\frac{\varrho'(d')}{d'}$ über jede Grenze oder kon-

Eckpunkte und, handelt es sich um einen Bezirk, der einen einzelnen Punkt bedeutet, diesen Punkt auch als Eckpunkt des Bezirks. In allen Fällen ist dann ein Eckpunkt p eines konvexen Bezirks $\mathfrak{R}$ dadurch charakterisiert, daß durch ihn drei Stützebenen an $\mathfrak{R}$ mit unabhängigen Normalenrichtungen gehen. Als *Eckstützebenen* durch p an $\mathfrak{R}$ bezeichnen wir jede solche Stützebene durch p an $\mathfrak{R}$, deren Bedingung sich auf irgendeine Weise in eine Form $c_1\varphi_1 + c_2\varphi_2 + c_3\varphi_3 \leq 0$ bringen läßt, so daß $\varphi_1 \leq 0$, $\varphi_2 \leq 0$, $\varphi_3 \leq 0$ die Bedingungen dreier Stützebenen an $\mathfrak{R}$ mit unabhängigen Normalenrichtungen und c_1, c_2, c_3 positive Konstanten sind. Für diese Eckstützebenen ist dann auch folgende Eigenschaft charakteristisch: Ist $(\alpha_0, \beta_0, \gamma_0)$ Normale einer Eckstützebene durch p an $\mathfrak{R}$, so ist eine jede Richtung (α, β, γ) , wobei die Beträge $|\alpha - \alpha_0|$, $|\beta - \beta_0|$, $|\gamma - \gamma_0|$ unter einer gewissen positiven Größe bleiben, ebenfalls stets Normale einer Stützebene durch p an $\mathfrak{R}$.

Wir denken uns zu jeder Bedingung

$$u(x - x_0) + v(y - y_0) + w(z - z_0) \leq 0$$

einer Stützebene durch p an $\mathfrak{R}$ den Punkt mit den Koordinaten u, v, w bestimmt und nehmen dazu noch den Nullpunkt $u = 0, v = 0, w = 0$ hinzu. Die dadurch hervorgehende, aus lauter Strahlen von o aus bestehende Punktmenge ist dann abgeschlossen, hat die Eigenschaft 1a) eines konvexen Bezirks und enthält den Nullpunkt o sowie drei nicht mit ihm in einer Ebene gelegene Punkte. Diese Eigenschaften 1c), 1a), 2) eines konvexen Körpers übertragen sich auch auf denjenigen Bereich, welchen diese Punktmenge mit der Kugel $\mathfrak{G}$ gemein hat und den wir mit $N(p)$ bezeichnen. Dieser Bereich $N(p)$ ist zudem ganz im Endlichen gelegen und wird nun tatsächlich einen konvexen Körper vorstellen. Dieser Körper $N(p)$ besteht aus allen Radien der Kugel $\mathfrak{G}$ von o nach den Punkten (α, β, γ) auf $\mathfrak{G}$, welche zugleich äußere Normalen von Stützebenen durch p an $\mathfrak{R}$ bedeuten, und wir wollen deshalb $N(p)$ den *Normalensektor* zum Eckpunkte p von $\mathfrak{R}$ nennen.

Enthält $N(p)$ den Punkt o als inneren Punkt, so muß $N(p)$ mit der ganzen Kugel $\mathfrak{G}$ zusammenfallen und also jede Stützebene an $\mathfrak{R}$ durch p laufen. Dieser Fall tritt nur ein, wenn der Bezirk $\mathfrak{R}$ den einen Punkt p bedeutet. In jedem anderen Falle ist o notwendig ein Punkt der Begrenzung von $N(p)$ und der Körper $N(p)$ zugleich sein eigener Projektionssektor von o aus.

Ist $\mathfrak{R}$ ein konvexer Körper und (α, β, γ) (äußere) Normale einer Stützebene durch p an $\mathfrak{R}$, so geht die Stützebene an $\mathfrak{R}$ mit der (äußeren) Normale $(-\alpha, -\beta, -\gamma)$ jedenfalls nicht durch p . Der Sektor $N(p)$ kann also hier nicht zwei Radien von o aus in entgegengesetzten Richtungen

enthalten. Der Punkt $\mathfrak{o}$ ist daher in diesem Falle ein Eckpunkt in $N(\mathfrak{p})$ und $N(\mathfrak{p})$ ein Kegelsektor. Bedeutet alsdann $(-\alpha^*, -\beta^*, -\gamma^*)$ die Normale einer Eckstützebene durch $\mathfrak{o}$ an $N(\mathfrak{p})$, so gilt für jede Normale (α, β, γ) einer Stützebene durch $\mathfrak{p}$ an $\mathfrak{R}$ stets die Beziehung

$$(61) \quad \alpha\alpha^* + \beta\beta^* + \gamma\gamma^* > 0.$$

Ist $\mathfrak{R}$ ein Oval, so ist die Ebene des Ovals und nur diese eine Ebene Stützebene an $\mathfrak{R}$ zugleich mit zwei einander entgegengesetzten (äußeren) Normalen. Alsdann enthält $N(\mathfrak{p})$ einen einzigen Durchmesser der Kugel $\mathfrak{G}$, der Punkt $\mathfrak{o}$ ist ein Kantenpunkt von $N(\mathfrak{p})$ und $N(\mathfrak{p})$ ein Keilsektor der Kugel $\mathfrak{G}$.

Wenn endlich $\mathfrak{R}$ eine Strecke bedeutet, so enthält $N(\mathfrak{p})$ gewiß zwei verschiedene Durchmesser von $\mathfrak{G}$, und ist $\mathfrak{o}$ ein Flächenpunkt von $N(\mathfrak{p})$ und daher $N(\mathfrak{p})$ eine Halbkugel von $\mathfrak{G}$.

Die *inneren* Radien von $N(\mathfrak{p})$, d. h. diejenigen, welche von $\mathfrak{o}$ aus ins Innere von $N(\mathfrak{p})$ eintreten, entsprechen den Normalen (α, β, γ) der *Eckstützebenen* durch $\mathfrak{p}$ an $\mathfrak{R}$. Da eine solche Eckstützebene von $\mathfrak{R}$ überhaupt nur den einen Punkt $\mathfrak{p}$ enthält, kann sie niemals durch einen zweiten Eckpunkt von $\mathfrak{R}$ gehen. Besitzt der konvexe Bezirk $\mathfrak{R}$ verschiedene Eckpunkte, so müssen danach die Normalensektoren zu diesen Eckpunkten untereinander in ihren inneren Punkten durchweg verschieden sein. Nun sind sie sämtlich Teile der Kugel $\mathfrak{G}$, deren Volumen endlich $= \frac{4\pi}{3}$ ist. Auch in dem Falle, daß die Zahl der Eckpunkte von $\mathfrak{R}$ eine unendliche ist, kann es daher unter diesen Eckpunkten jedesmal nur eine endliche Anzahl geben, bei welchen das Volumen des zugehörigen Normalensektors eine beliebig angenommene positive Größe übersteigt. Wir sind danach stets imstande, die sämtlichen Eckpunkte eines konvexen Bezirks nach dem Volumen ihrer Normalensektoren in eine abzählbare Reihe zu ordnen.

Ist $\mathfrak{R}$ ein konvexer Bezirk mit einer endlichen Anzahl von extremen Punkten, also ein Polyeder, Polygon, eine Strecke oder ein Punkt, so sind alle extremen Punkte von $\mathfrak{R}$ zugleich Eckpunkte von $\mathfrak{R}$. Es geht daher jede Stützebene an $\mathfrak{R}$ durch wenigstens einen Eckpunkt von $\mathfrak{R}$, und gehört infolgedessen jeder Radius von $\mathfrak{G}$ dem Normalensektor zu wenigstens einem Eckpunkte von $\mathfrak{R}$ an. Die Normalensektoren der sämtlichen Eckpunkte von $\mathfrak{R}$ erfüllen dann also die ganze Kugel $\mathfrak{G}$ und die Summe ihrer Volumina ist genau $= \frac{4\pi}{3}$.

§ 16. Normalensektor zu einem äußeren Punkte eines konvexen Körpers. Kappenkörper.

Es sei jetzt $\mathfrak{K}$ ein konvexer Körper und p mit den Koordinaten x_0, y_0, z_0 ein beliebiger Punkt außerhalb $\mathfrak{K}$. Es bedeute $H(u, v, w)$ die Stützebenenfunktion von $\mathfrak{K}$, diejenige des Punktes p ist

$$(62) \quad H_0(u, v, w) = ux_0 + vy_0 + wz_0.$$

Bezeichnen wir das Maximum unter den zwei Werten $H(u, v, w)$ und $H_0(u, v, w)$ jedesmal mit $H^*(u, v, w)$, so ist $H^*(u, v, w)$ nach § 11 die Stützebenenfunktion des kleinsten, den Körper $\mathfrak{K}$ und den Punkt p zusammen enthaltenden konvexen Körpers $(p, \mathfrak{K})$. Die Menge aller derjenigen Punkte, welche dieser Körper $(p, \mathfrak{K})$ außer den Punkten von $\mathfrak{K}$ enthält, möge die *Kappe* von p an $\mathfrak{K}$ heißen.

Wir wollen die Punktmenge aller Strecken $p\mathfrak{f}$ von p nach den verschiedenen Punkten $\mathfrak{f}$ des Körpers $\mathfrak{K}$ mit $\mathfrak{K}^*$ bezeichnen. Wir stellen nun vor allem fest, daß der Körper $(p, \mathfrak{K})$ sich genau mit dieser Punktmenge $\mathfrak{K}^*$ deckt. In der Tat, zunächst enthält $(p, \mathfrak{K})$ als konvexer Bezirk jede von jenen Strecken $p\mathfrak{f}$ und also die ganze Menge $\mathfrak{K}^*$.

Diese Menge $\mathfrak{K}^*$ aber bildet bereits für sich einen konvexen Bezirk. Sie besitzt offenbar die Eigenschaft 1a) eines konvexen Bezirks deshalb, weil diese Eigenschaft dem Körper $\mathfrak{K}$ zukommt; zweitens ist sie wie $\mathfrak{K}$ und p ganz im Endlichen gelegen. Drittens ist sie nun auch eine abgeschlossene Punktmenge. Denn erscheint irgendein Punkt r im Raume als Häufungsstelle von unendlich vielen Punkten der Form $(1-t)p + t\mathfrak{f}$, wobei jedesmal $\mathfrak{f}$ einen Punkt von $\mathfrak{K}$ und t einen Wert ≥ 0 und ≤ 1 bedeutet, so können wir aus diesen Punkten eine unendliche Reihe solcher herausnehmen, die nach dem Punkte r konvergieren, d. h. deren Abstände von r nach Null abnehmen. Da $\mathfrak{K}$ ganz in einer Kugel von endlichem Radius liegt, unendlich viele Punkte $\mathfrak{f}$ aus $\mathfrak{K}$ also stets wenigstens eine Häufungsstelle besitzen müssen, können wir dann weiter aus jener ersten Reihe von Punkten eine zweite unendliche Reihe

$$(1-t_1)p + t_1\mathfrak{f}_1, (1-t_2)p + t_2\mathfrak{f}_2, \dots$$

aussondern, für die auch noch die darin vorkommende Punktreihe $\mathfrak{f}_1, \mathfrak{f}_2, \dots$ nach einem Grenzpunkte konvergiert. Da $\mathfrak{K}$ eine abgeschlossene Punktmenge vorstellt, ist dieser Grenzpunkt wieder irgendein Punkt $\mathfrak{f}$ aus $\mathfrak{K}$, und gleichzeitig muß auch die Reihe der Werte $t_1, t_2, \dots$ nach einem Grenzwerte $t \geq 0$ und ≤ 1 konvergieren, mit dem dann $r = (1-t)p + t\mathfrak{f}$ gilt. Danach liegt die Häufungsstelle r der Menge $\mathfrak{K}^*$ selbst auf einer Strecke von p nach einem Punkte von $\mathfrak{K}$, gehört somit selbst zu $\mathfrak{K}^*$, d. h. eben $\mathfrak{K}^*$ ist eine abgeschlossene Punktmenge; $\mathfrak{K}^*$ besitzt also in der Tat

alle Eigenschaften eines konvexen Bezirks. Da nun der Körper $(p, \mathfrak{R})$ in jedem, p und $\mathfrak{R}$ enthaltenden konvexen Bezirk enthalten ist, muß danach umgekehrt $(p, \mathfrak{R})$ in $\mathfrak{R}^*$ enthalten sein und deckt sich also genau mit der Menge $\mathfrak{R}^*$.

Betrachten wir weiter alle Punkte der Form $p + t(f - p)$, wobei f einen Punkt von $\mathfrak{R}$ und t jetzt einen beliebigen Wert ≥ 0 bedeute, so ist die Menge $\mathfrak{R}$ dieser Punkte ebenso wie $\mathfrak{R}$ abgeschlossen und stellt den Projektionsraum des Körpers $(p, \mathfrak{R})$ von p aus vor. Durch p können wir jedenfalls eine solche Ebene legen, welche $\mathfrak{R}$ nicht trifft, und diese enthält dann vom Raume $\mathfrak{R}$ einzig den Punkt p , danach ist p ein *Eckpunkt* im Körper $(p, \mathfrak{R})$.

Als (äußere) Normalen der Stützebenen durch p an $(p, \mathfrak{R})$ haben wir diejenigen Richtungen (α, β, γ) , bei welchen

$$(63) \quad H^*(\alpha, \beta, \gamma) = \alpha x_0 + \beta y_0 + \gamma z_0 = H_0(\alpha, \beta, \gamma)$$

oder also, was damit gleichbedeutend ist,

$$(64) \quad H(\alpha, \beta, \gamma) \leq \alpha x_0 + \beta y_0 + \gamma z_0$$

gilt. Die Radien der Kugel $\mathfrak{G}$, die von o aus in diesen Richtungen führen, bilden den Normalensektor zur Ecke p im Körper $(p, \mathfrak{R})$, den wir jetzt den *Normalensektor zum äußeren Punkte p von $\mathfrak{R}$* nennen und wieder mit $N(p)$ bezeichnen. Die *inneren* Radien von $N(p)$ bestimmen uns die Normalen (α, β, γ) der *Eckstützebenen* durch p an $(p, \mathfrak{R})$; für diese Richtungen gilt

$$(65) \quad H(\alpha, \beta, \gamma) < \alpha x_0 + \beta y_0 + \gamma z_0,$$

sie bedeuten also gleichzeitig die Normalen derjenigen Stützebenen an $\mathfrak{R}$, welche p von $\mathfrak{R}$ trennen, d. h. p auf entgegengesetzter Seite wie das Innere von $\mathfrak{R}$ liegen haben.

Es sei sodann $\bar{N}(p)$ der zu $N(p)$ komplementäre abgeschlossene Sektor in der Kugel, d. h. die Punktmenge aller Radien von o nach solchen Punkten (α, β, γ) auf $\mathfrak{G}$, für welche

$$(66) \quad H(\alpha, \beta, \gamma) \geq \alpha x_0 + \beta y_0 + \gamma z_0$$

oder also

$$(67) \quad H^*(\alpha, \beta, \gamma) = H(\alpha, \beta, \gamma)$$

ist. Die beiden Sektoren $N(p)$ und $\bar{N}(p)$ bestimmen einander gegenseitig, indem sie zusammen die ganze Kugel $\mathfrak{G}$ ausfüllen, in ihren inneren Punkten durchweg verschieden sind, ihre begrenzenden Radien dagegen völlig gemeinsam haben.

Entsprechend seiner Stützebenenfunktion $H^*(u, v, w)$ hat der Körper $\mathfrak{R}^* = (p, \mathfrak{R})$ jede solche Stützebene, deren Normale der Bedingung (67), also einem Radius von $\bar{N}(p)$ entspricht, mit dem Körper $\mathfrak{R}$ gemein; es

erfüllt also $(p, \mathfrak{K})$ die Bedingungen

$$(68) \quad \alpha x + \beta y + \gamma z \leq H(\alpha, \beta, \gamma)$$

für alle Richtungen (α, β, γ) , die den Radien von $\bar{N}(p)$ entsprechen. Von den bezüglichlichen Stützebenen (68) gehen diejenigen, welche den begrenzenden Radien von $\bar{N}(p)$ entsprechen, durch den Punkt p . Da $\bar{N}(p)$ und die inneren Radien von $N(p)$ die ganze Kugel $\mathfrak{G}$ zusammensetzen, besitzt der Körper $(p, \mathfrak{K})$ an weiteren Stützebenen nur die Eckstützebenen durch p , welche von diesem Körper einzig den Punkt p enthalten. Hiernach nehmen bereits die zuerst genannten Stützebenen die ganze Begrenzung von $(p, \mathfrak{K})$ in sich auf, und ist daher der Körper $(p, \mathfrak{K})$ durch die Gesamtheit der Bedingungen (68), (66) schon vollkommen charakterisiert. Auf diese Weise ist durch $\mathfrak{K}$ und den in der Kugel $\mathfrak{G}$ konstruierten Sektor $\bar{N}(p)$ bzw. den Sektor $N(p)$ der Körper $(p, \mathfrak{K})$ und sodann der Punkt p als der einzige, nicht zu $\mathfrak{K}$ gehörende Eckpunkt von $(p, \mathfrak{K})$ genau bestimmt.

Es seien jetzt p und q zwei verschiedene Punkte außerhalb $\mathfrak{K}$, so können wir das Verhalten der durch p und q gelegten Geraden zum Körper $\mathfrak{K}$ mit Hilfe der zu p und q gehörigen Normalensektoren $N(p)$ und $N(q)$ in einfacher Weise beurteilen. Wir haben folgende Möglichkeiten zu erwägen:

1. Die Gerade pq treffe den Körper $\mathfrak{K}$ nicht auf der Strecke pq , aber auf deren Verlängerung über q hinaus. Alsdann ist q ein Punkt der Kappe von p an $\mathfrak{K}$ und daher der Körper $(q, \mathfrak{K})$ ganz in $(p, \mathfrak{K})$ und weiter die Kappe von q an $\mathfrak{K}$ ganz in der Kappe von p an $\mathfrak{K}$ enthalten. Andererseits sind dann diejenigen Stützebenen an $\mathfrak{K}$, welche gleichzeitig Stützebenen an $(p, \mathfrak{K})$ vorstellen, notwendig auch Stützebenen an $(q, \mathfrak{K})$ und ist daher der Sektor $\bar{N}(p)$ ganz in dem Sektor $\bar{N}(q)$ enthalten, also enthält endlich der komplementäre Sektor $N(p)$ ganz den komplementären Sektor $N(q)$.

Umgekehrt erkennen wir, daß, wenn der Sektor $N(p)$ den Sektor $N(q)$ enthält, notwendig für die Gerade pq der hier angenommene Fall vorliegt. Denn es ist dann andererseits der komplementäre Sektor $\bar{N}(p)$ ganz in dem komplementären Sektor $\bar{N}(q)$ enthalten, und genügt daher der Körper $(q, \mathfrak{K})$ allen den Ungleichungen (68), welche in ihrer Gesamtheit genau den Bereich des Körpers $(p, \mathfrak{K})$ charakterisieren. Mithin ist q selbst in $(p, \mathfrak{K})$ enthalten, also ein Punkt der Kappe von p an $\mathfrak{K}$.

2. Die Gerade pq treffe den Körper $\mathfrak{K}$ nicht auf der Strecke pq , aber auf deren Verlängerung über p hinaus. Für diesen Fall ist charakteristisch, daß der Sektor $N(p)$ ganz in dem Sektor $N(q)$ enthalten ist.

3. Die Gerade pq treffe den Körper $\mathfrak{K}$ überhaupt nicht. Der Punkt q gehört dann weder dem Projektionsraume $\mathfrak{N}$ des Körpers $(p, \mathfrak{K})$ von p

aus noch dem dazu in bezug auf p symmetrischen Kegelraume an, und gibt es daher nach § 14 Ebenen durch pq , welche mit $\mathfrak{R}$ nur den Punkt p gemein haben, also keinen Punkt von $\mathfrak{R}$ enthalten, mithin gleichzeitig Eckstützebenen durch p an $(p, \mathfrak{R})$ und durch q an $(q, \mathfrak{R})$ sind. In diesem Falle haben daher die Normalensektoren $N(p)$ und $N(q)$ innere Punkte gemein, es ist aber nach dem bei 1. und 2. Ausgeführten keiner dieser Sektoren vollständig in dem anderen enthalten.

4. Endlich enthalte die Strecke pq selbst wenigstens einen Punkt t von $\mathfrak{R}$. Alsdann gibt es keine Stützebene an $\mathfrak{R}$, welche p und q beide auf der entgegengesetzten Seite wie das Innere von $\mathfrak{R}$ liegen hat. In diesem Falle haben daher die Normalensektoren $N(p)$ und $N(q)$ keinen inneren Punkt gemein, und umgekehrt erkennen wir letzteren Umstand als charakteristisch für die hier angenommene Lage der Punkte p und q zu $\mathfrak{R}$, wenn wir damit die Ergebnisse in den zuerst untersuchten Fällen vergleichen. Ist dann p^* ein Punkt der Kappe von p an $\mathfrak{R}$, q^* ein Punkt der Kappe von q an $\mathfrak{R}$, so ist der Normalensektor $N(p^*)$ im Sektor $N(p)$ und der Normalensektor $N(q^*)$ im Sektor $N(q)$ enthalten und haben daher auch $N(p^*)$ und $N(q^*)$ keinen inneren Punkt gemein. Somit kann niemals p^* mit q^* identisch sein, und zudem enthält die Strecke p^*q^* stets wenigstens einen Punkt t^* von $\mathfrak{R}$ und läßt sich dann in zwei Strecken p^*t^* und t^*q^* zerlegen, von denen die eine ganz dem Körper $(p, \mathfrak{R})$, die andere ganz dem Körper $(q, \mathfrak{R})$ angehört.

Auf Grund der letzten Bemerkung leiten wir noch folgenden Satz her:

Es sei $\mathfrak{R}$ ein konvexer Körper und es seien $p_1, p_2, \dots$ eine endliche oder unendliche Reihe von Punkten außerhalb $\mathfrak{R}$ derart, daß von ihren Normalensektoren $N(p_1), N(p_2), \dots$ in bezug auf $\mathfrak{R}$ keine zwei untereinander einen inneren Punkt gemein haben. Alsdann bildet die Vereinigung der sämtlichen Körper $(p_1, \mathfrak{R}), (p_2, \mathfrak{R}), \dots$ (d. i. also der Körper $\mathfrak{R}$ unter Hinzunahme der Kappen von allen einzelnen Punkten $p_1, p_2, \dots$ an $\mathfrak{R}$) stets wieder einen konvexen Körper.

Wir bezeichnen den in dieser Art gebildeten Bereich mit

$$\mathfrak{R}^* = (p_1, p_2, \dots, \mathfrak{R}).$$

Aus der zuletzt gemachten Bemerkung erhellt, daß dieser Bereich jedenfalls die Eigenschaft 1a) eines konvexen Bezirks hat, mit irgend zwei Punkten p^*, q^* stets die ganze Strecke p^*q^* zu enthalten. Ist die vorausgesetzte Punktreihe $p_1, p_2, \dots$ eine endliche, so versteht sich ferner, daß ebenso wie jeder einzelne der Körper $(p_1, \mathfrak{R}), (p_2, \mathfrak{R}), \dots$ auch dieser ganze Bereich eine abgeschlossene Punktmenge und ganz im Endlichen gelegen ist. In diesem Falle bedeutet dann $\mathfrak{R}^*$ einfach den kleinsten, $\mathfrak{R}$ und die sämtlichen Punkte $p_1, p_2, \dots$ aufnehmenden konvexen Bezirk.

Stellt $\mathfrak{K}$ ein Polyeder vor, so ist der soeben angenommene Fall, daß die Reihe $p_1, p_2, \dots$ eine endliche ist, überhaupt allein möglich. Denn ist dann p ein Punkt außerhalb $\mathfrak{K}$, so muß es unter den Ebenen der Seitenflächen des Polyeders $\mathfrak{K}$ wenigstens eine geben, welche p auf entgegengesetzter Seite liegen hat wie das Innere von $\mathfrak{K}$. Der Radius der Kugel $\mathfrak{G}$ in Richtung der Normale (α, β, γ) dieser Seitenfläche tritt dann als ein innerer Radius in dem zu p gehörenden Normalensektor $N(p)$ auf. Danach kann die vorausgesetzte Reihe von Punkten $p_1, p_2, \dots$ außerhalb $\mathfrak{K}$ mit Normalensektoren $N(p_1), N(p_2), \dots$, die im Inneren durchweg verschieden sind, gewiß nicht mehr Punkte aufweisen, als das Polyeder $\mathfrak{K}$ Seitenflächen besitzt.

Es sei jetzt der konvexe Körper $\mathfrak{K}$ kein Polyeder und die vorausgesetzte Reihe der Punkte $p_1, p_2, \dots$ eine unendliche. Wir wollen uns der Einfachheit wegen den Nullpunkt o im Inneren von $\mathfrak{K}$ gelegen denken. Ist ε eine beliebige positive Größe, so können wir dann nach § 5 zu $\mathfrak{K}$ stets ein Polyeder $\mathfrak{Q}$ bestimmen, welches den Körper $\mathfrak{K}$ enthält und selbst in $(1 + \varepsilon)\mathfrak{K}$ enthalten ist. Ist p ein Punkt außerhalb $\mathfrak{Q}$, so ist jede Stützebene durch p an den Körper $(p, \mathfrak{Q})$ zugleich eine Stützebene durch p an $(p, \mathfrak{K})$ und ist daher der Normalensektor zum Eckpunkte p von $(p, \mathfrak{Q})$ ganz im Normalensektor $N(p)$ zum Eckpunkte p von $(p, \mathfrak{K})$ enthalten. Haben wir verschiedene Punkte p außerhalb $\mathfrak{Q}$, deren Normalensektoren in bezug auf $\mathfrak{K}$ in ihren inneren Punkten verschieden sind, so werden daher auch gewiß deren Normalensektoren in bezug auf $\mathfrak{Q}$ in ihren inneren Punkten verschieden sein. Da $\mathfrak{Q}$ ein Polyeder ist, kann es nach dem vorhin Bemerkten infolgedessen in der Reihe $p_1, p_2, \dots$ jedesmal nur eine endliche Anzahl von Punkten geben, die außerhalb $\mathfrak{Q}$, und um so mehr nur eine endliche Anzahl geben, die außerhalb $(1 + \varepsilon)\mathfrak{K}$ liegen. Daraus entnehmen wir weiter, daß jede irgend vorhandene Häufungsstelle der Reihe $p_1, p_2, \dots$ notwendig ein Punkt auf der Begrenzung von $\mathfrak{K}$ ist, und daß auch in diesem Falle, wo es sich um unendlich viele Punkte $p_1, p_2, \dots$ handelt, die Körper $(p_1, \mathfrak{K}), (p_2, \mathfrak{K}), \dots$ zusammen in einer Kugel von endlichem Radius enthalten sind und ihre Vereinigung $\mathfrak{K}^*$ eine abgeschlossene Punktmenge vorstellt, somit in der Tat $\mathfrak{K}^*$ wieder alle Eigenschaften eines konvexen Körpers besitzt.

Der Körper $\mathfrak{K}^*$ ist in jedem Falle vollständig charakterisiert als der kleinste konvexe Körper, der $\mathfrak{K}$ und die sämtlichen Punkte $p_1, p_2, \dots$ aufnimmt; wir bezeichnen deshalb $\mathfrak{K}^*$ auch mit $(p_1, p_2, \dots, \mathfrak{K})$ und wir wollen jeden in dieser Art aus $\mathfrak{K}$ abgeleiteten konvexen Körper einen *Kappenkörper* von $\mathfrak{K}$ nennen. Die Punkte p_h sind sämtlich Eckpunkte in $\mathfrak{K}^*$, da $\mathfrak{K}^*$ ganz im Projektionsraume des Körpers $(p_h, \mathfrak{K})$ von p_h aus liegt. Jede Stützebene an $\mathfrak{K}^*$ ist notwendig Stützebene an wenigstens

einen der Körper $(p_h, \mathfrak{R})$, also entweder Stützebene an $\mathfrak{R}$ oder Eckstützebene an $\mathfrak{R}^*$ durch einen Eckpunkt p_h . Umgekehrt gilt folgende Tatsache:

Sind für einen konvexen Körper $\mathfrak{R}^*$ alle Stützebenen mit Ausnahme der Eckstützebenen durch gewisse Eckpunkte $p_1, p_2, \dots$ zugleich Stützebenen an den konvexen Körper $\mathfrak{R}$, so ist $\mathfrak{R}^*$ ein Kappenkörper von $\mathfrak{R}$.

Denn da in einem Eckpunkte p_h von $\mathfrak{R}^*$ die Bedingungen der Eckstützebenen aus den Bedingungen der anderen Stützebenen durch p_h an $\mathfrak{R}^*$ folgen, so erfüllt $\mathfrak{R}$ jedenfalls die Bedingungen aller Stützebenen an $\mathfrak{R}^*$, es ist also $\mathfrak{R}$ in $\mathfrak{R}^*$ enthalten und liegen $p_1, p_2, \dots$ außerhalb $\mathfrak{R}$. Sodann sind die Normalensektoren zu den Eckpunkten $p_1, p_2, \dots$ von $\mathfrak{R}^*$ untereinander im Inneren verschieden, und gilt daher das Nämliche von den Normalensektoren zu $p_1, p_2, \dots$ in bezug auf $\mathfrak{R}$, so daß wir den Kappenkörper $(p_1, p_2, \dots, \mathfrak{R})$ von $\mathfrak{R}$ herstellen können. Jede Stützebene an $\mathfrak{R}^*$ ist dann zugleich Stützebene an diesen Kappenkörper und ist daher $\mathfrak{R}^*$ mit letzterem Körper identisch.

Nehmen wir für den Körper $\mathfrak{R}$ speziell die Kugel $\mathfrak{G}$ mit o als Mittelpunkt vom Radius 1 und ist p ein Punkt außerhalb $\mathfrak{G}$, so ergänzen die Kappe von p an $\mathfrak{G}$ und der Normalensektor $N(p)$ dieser Kappe sich genau zu dem Doppelkegel, der durch Rotation eines Dreiecks otp , wobei die Länge $ot = 1$ und der Winkel otp ein rechter ist, um op entsteht. Der allgemeinste Kappenkörper der Kugel $\mathfrak{G}$ wird sodann erhalten, indem man auf der Oberfläche dieser Kugel eine endliche oder unendliche Reihe von Kalotten bezeichnet, von denen eine jede kleiner als eine Halbkugel ist und die untereinander in ihren inwendigen Punkten durchweg verschieden sind, und auf jede Kalotte den Rotationskegel aufsetzt, der die Kalotte als Basis hat und dessen Spitze in solcher Distanz von o liegt, daß die Strecken auf dem Mantel die Kugel berühren.

III. Kapitel.

Scharen konvexer Körper.

§ 17. Schar konvexer Bezirke.

Bedeutend $p_1, p_2, \dots, p_m$ irgendwelche Punkte und sind x_i, y_i, z_i die Werte der Koordinaten x, y, z von p_i (für $i = 1, 2, \dots, m$), sind ferner $s_1, s_2, \dots, s_m$ irgend m Konstanten, so soll nach einer in § 4 getroffenen Festsetzung unter

$$(69) \quad p = s_1 p_1 + s_2 p_2 + \dots + s_m p_m$$

derjenige Punkt verstanden werden, dessen Koordinaten durch

$$(70) \quad x = \sum s_i x_i, \quad y = \sum s_i y_i, \quad z = \sum s_i z_i \quad (i = 1, 2, \dots, m)$$

bestimmt sind. Die Bedeutung dieses Punktes p ist stets dann völlig unabhängig von dem zugrunde gelegten Systeme von Parallelkoordinaten x, y, z , wenn $s_1 + s_2 + \dots + s_m = 1$ ist, dagegen, wenn nicht die letztere Beziehung statthat, noch wesentlich abhängig von dem als Anfangspunkt der Koordinaten angenommenen Punkte o , wie die Darstellung desselben Punktes in der Form

$$(71) \quad p = s_1 p_1 + s_2 p_2 + \dots + s_m p_m + (1 - s_1 - s_2 - \dots - s_m) o$$

zum Ausdrucke bringt.

Es seien $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ eine endliche Reihe von konvexen Bezirken, von denen also ein jeder einen konvexen Körper oder ein ebenes Oval oder eine Strecke oder einen Punkt bedeuten kann, und

$$(72) \quad H_1(u, v, w), H_2(u, v, w), \dots, H_m(u, v, w)$$

ihre Stützebenenfunktionen. Jede dieser Funktionen genügt den Bedingungen 1., 2., 3. in § 11. Sind nun $s_1, s_2, \dots, s_m$ beliebige solche konstante Werte, die sämtlich ≥ 0 sind, so können wir aus jenen Bedingungen für diese einzelnen Funktionen sofort die entsprechenden Beziehungen für die Funktion

$$(73) \quad H(u, v, w) = s_1 H_1(u, v, w) + s_2 H_2(u, v, w) + \dots + s_m H_m(u, v, w)$$

herleiten, und wird daher nach einem Satze in § 11 diese lineare Verbindung $H = \sum s_i H_i$ jedesmal wieder die Stützebenenfunktion eines gewissen konvexen Bezirks $\mathfrak{R}$ vorstellen. Für diesen Bezirk $\mathfrak{R}$ können wir sodann die folgende zweite Entstehungsweise aus den Bezirken $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ feststellen:

Der konvexe Bezirk $\mathfrak{R}$ mit der Stützebenenfunktion

$$H = \sum s_i H_i \quad (i = 1, 2, \dots, m)$$

ist identisch mit der Menge aller derjenigen Punkte p , welche sich auf irgendeine Weise in die Form

$$p = \sum s_i p_i \quad (i = 1, 2, \dots, m)$$

setzen lassen, so daß dabei jedesmal p_i ein Punkt aus dem konvexen Bezirke $\mathfrak{R}_i$ ist.

In der Tat, bezeichnen wir vorläufig die Menge aller Punkte p , welche irgendwie in der eben erwähnten Form $\sum s_i p_i$ darzustellen sind, mit $\mathfrak{R}^*$, so leuchtet zunächst ein, daß der konvexe Bezirk $\mathfrak{R}$ mit der Stützebenenfunktion $H = \sum s_i H_i$ jedenfalls diese Punktmenge $\mathfrak{R}^*$ ganz enthält. Denn sind u, v, w irgendwelche Werte, so gilt für einen Punkt p_i mit den Koordinaten x_i, y_i, z_i aus $\mathfrak{R}_i$ jedesmal

$$(74) \quad ux_i + vy_i + wz_i \leq H_i(u, v, w),$$

und infolgedessen für die Koordinaten x, y, z des Punktes $p = \sum s_i p_i$ zufolge (70) und (73) dann

$$(75) \quad ux + vy + wz \leq H(u, v, w),$$

so daß ein beliebiger derartiger Punkt p aus $\mathfrak{R}^*$ notwendig stets zum Bezirke $\mathfrak{R}$ mit der Stützebenenfunktion $H = \sum s_i H_i$ gehört.

Ferner können wir zeigen, daß die Punktmenge $\mathfrak{R}^*$ für sich bereits einen konvexen Bezirk vorstellt. Haben wir für zwei verschiedene Punkte p' und p'' aus $\mathfrak{R}^*$ Darstellungen

$$(76) \quad p' = \sum s_i p'_i, \quad p'' = \sum s_i p''_i \quad (i = 1, 2, \dots, m),$$

wo p'_i und p''_i Punkte aus $\mathfrak{R}_i$ sind, und ist t ein Wert > 0 und < 1 , so folgt

$$(77) \quad (1-t)p' + tp'' = \sum s_i ((1-t)p'_i + tp''_i),$$

und ist daraus nach der Eigenschaft 1a) für die Bezirke $\mathfrak{R}_i$ sofort ersichtlich, daß auch jeder Punkt der Strecke $p'p''$ als Punkt der Menge $\mathfrak{R}^*$ auftritt, somit dieser Menge $\mathfrak{R}^*$ ebenfalls die Eigenschaft 1a) eines konvexen Bezirks zukommt. Weiter bestehen für die Koordinaten der Punkte p_i in $\mathfrak{R}_i$ untere und obere Grenzen und erschließen wir daraus mit Rücksicht auf (70) sogleich eine untere und obere Grenze für die Koordinaten der Punkte in $\mathfrak{R}^*$, so daß $\mathfrak{R}^*$ die Eigenschaft 1b) eines konvexen Bezirks hat. Endlich ist $\mathfrak{R}^*$ auch eine abgeschlossene Punktmenge. Denn ist irgendein Punkt q eine Häufungsstelle von $\mathfrak{R}^*$, so können wir in $\mathfrak{R}^*$ eine unendliche Reihe von Punkten

$$(78) \quad p^{(1)}, p^{(2)}, p^{(3)}, \dots$$

angeben, welche nach q konvergieren, d. h. deren Abstände von q mit wachsendem Index nach Null abnehmen. Jeden dieser Punkte $p^{(x)}$ können wir auf irgendeine Art in die Form

$$(79) \quad p^{(x)} = s_1 p_1^{(x)} + s_2 p_2^{(x)} + \dots + s_m p_m^{(x)}$$

setzen, so daß dabei $p_i^{(x)}$ ein Punkt aus $\mathfrak{R}_i$ ist. Nun können wir, da $\mathfrak{R}_1$ ganz im Endlichen liegt, zunächst aus der Reihe $p_1^{(1)}, p_1^{(2)}, p_1^{(3)}, \dots$ eine solche unendliche Reihe $p_1^{(x_1')}, p_1^{(x_2')}, p_1^{(x_3')}, \dots$ aussondern, die nach einem bestimmten Grenzpunkte konvergiert, der p_1 heiße, sodann können wir, da $\mathfrak{R}_2$ ganz im Endlichen liegt, aus der Reihe $p_2^{(x_1')}, p_2^{(x_2')}, p_2^{(x_3')}, \dots$ eine solche unendliche Reihe $p_2^{(x_1'')}, p_2^{(x_2'')}, p_2^{(x_3'')}, \dots$ aussondern, die nach einem bestimmten Grenzpunkte p_2 konvergiert, usw. Zuletzt kommen wir dann zu einer unendlichen Reihe von Indizes $x_1^{(m)}, x_2^{(m)}, x_3^{(m)}, \dots$ aus der Reihe $1, 2, 3, \dots$ derart, daß zu gleicher Zeit für jeden Index $i = 1, 2, \dots, m$ die Reihe

$$(80) \quad p_i^{(x_1^{(m)})}, p_i^{(x_2^{(m)})}, p_i^{(x_3^{(m)})}, \dots$$

je nach einem bestimmten Grenzpunkte p_i konvergiert. Dabei entspringt für q als Grenzpunkt der Reihe (78) notwendig die Darstellung

$$(81) \quad q = s_1 p_1 + s_2 p_2 + \cdots + s_m p_m.$$

Nun gehört jedesmal der Grenzpunkt p_i zu $\mathfrak{R}_i$ selbst, da diese Menge abgeschlossen ist. Durch den letzten Ausdruck erscheint daher die Häufungsstelle q von $\mathfrak{R}^*$ selbst stets als ein Punkt dieser Menge $\mathfrak{R}^*$. Hiernach besitzt $\mathfrak{R}^*$ auch die dritte, für einen konvexen Bezirk zu fordernde Eigenschaft, abgeschlossen zu sein.

Es sei jetzt $H^*(u, v, w)$ die Stützebenenfunktion von $\mathfrak{R}^*$, so gilt einerseits stets $H^*(u, v, w) \leq H(u, v, w)$, weil $\mathfrak{R}^*$ in $\mathfrak{R}$ enthalten ist. Wir können aber, wenn u, v, w feste Werte sind, dazu für p_i jedesmal einen solchen Punkt x_i, y_i, z_i aus $\mathfrak{R}_i$ wählen, daß für ihn genau

$$ux_i + vy_i + wz_i = H_i(u, v, w)$$

ist, und dann wird für die Koordinaten x, y, z von $p = \sum s_i p_i$ genau

$$ux + vy + wz = H(u, v, w)$$

sein. Also ist das Maximum von $ux + vy + wz$ im Bezirke $\mathfrak{R}^*$, nämlich $H^*(u, v, w)$, andererseits $\geq H(u, v, w)$; es folgt danach allgemein $H^*(u, v, w) = H(u, v, w)$, $\mathfrak{R}^* = \mathfrak{R}$, was zu beweisen war.

Im Hinblick auf den eben bewiesenen Satz bezeichnen wir den konvexen Bezirk mit der Stützebenenfunktion $H = \sum s_i H_i$ auch schlechthin durch $\mathfrak{R} = \sum s_i \mathfrak{R}_i$. Die Gesamtheit der konvexen Bezirke $\sum s_i \mathfrak{R}_i$ für alle möglichen Parameterwerte $s_1, s_2, \dots, s_m$, die sämtlich ≥ 0 sind, nennen wir die *Schar konvexer Bezirke mit den Grundbezirken* $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$. Diejenigen Bezirke der Schar, welche mit lauter Parameterwerten $s_1, s_2, \dots, s_m$, die > 0 sind, konstruiert werden, sollen *inwendige* Bezirke der Schar, die anderen Bezirke, für welche ein Teil der Parameter $s_1, s_2, \dots, s_m$ oder diese sämtlich $= 0$ sind, *Randbezirke* der Schar heißen. Es ist dabei nicht ausgeschlossen, daß ein und derselbe konvexe Bezirk in der Schar mittels verschiedener Wertsysteme der Parameter und insbesondere sowohl als ein inwendiger wie als ein Randbezirk auftritt.

Sind $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ speziell m ebene Bezirke (Ovale, Strecken, Punkte) in m parallelen Ebenen, welche die Richtung (α, β, γ) als gemeinsame Normale haben, so gilt bei jedem Index $i = 1, 2, \dots, m$ stets $H_i(\alpha, \beta, \gamma) + H_i(-\alpha, -\beta, -\gamma) = 0$ und daher auch für jeden Bezirk $\mathfrak{R} = \sum s_i \mathfrak{R}_i$ stets $H(\alpha, \beta, \gamma) + H(-\alpha, -\beta, -\gamma) = 0$. Es ist dann also jeder einzelne Bezirk der betreffenden, aus $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ entspringenden Schar ganz in einer Ebene mit der Normale (α, β, γ) gelegen.

Gibt es keine Richtung (α, β, γ) , wofür alle m Gleichungen

$$(82) \quad H_i(\alpha, \beta, \gamma) + H_i(-\alpha, -\beta, -\gamma) = 0 \quad (i = 1, 2, \dots, m)$$

auf einmal gelten, sind also $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ nicht m ebene Bezirke in m parallelen Ebenen, so wird für einen inwendigen Bezirk $\mathfrak{R}$ der aus ihnen entspringenden Schar notwendig bei jeder beliebigen Richtung (α, β, γ) immer

$$(83) \quad H(\alpha, \beta, \gamma) + H(-\alpha, -\beta, -\gamma) > 0$$

sein. In diesem Falle sind daher die *inwendigen* Bezirke der Schar durchweg konvexe Körper und soll die Schar als *Schar konvexer Körper* bezeichnet werden; dabei wird zugelassen, daß als Randbezirke der Schar auch Ovale, Strecken, Punkte auftreten.

§ 18. Die Begrenzungen in den Bezirken einer Schar.

Für die soeben festgestellte Bildungsweise der Bezirke in einer Schar aus den Grundbezirken geben wir jetzt noch eine zweite Ableitung, durch welche wir zugleich wichtige Aufschlüsse über die Begrenzung eines beliebigen Bezirks der Schar erhalten.

Es seien $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ irgend m konvexe Bezirke mit den Stützebenenfunktionen $H_1, H_2, \dots, H_m$ und $\mathfrak{R}$ der Bezirk mit der Stützebenenfunktion $H = \sum s_i H_i$ ($i = 1, 2, \dots, m$), wobei wir jetzt die Werte $s_1, s_2, \dots, s_m$ sämtlich > 0 voraussetzen.

Betrachten wir einen beliebigen extremen Punkt p von $\mathfrak{R}$; es sei S :

$$\xi = \alpha''x + \beta''y + \gamma''z, \quad \eta = \alpha'x + \beta'y + \gamma'z, \quad \xi = \alpha x + \beta y + \gamma z$$

eine orthogonale Substitution, zu der dieser extreme Punkt von $\mathfrak{R}$ gehört. Wir bezeichnen wie in § 12 mit $\{\mathfrak{F}\}$ die Stützebene an $\mathfrak{R}$ mit der Normale (α, β, γ) , mit $\mathfrak{F}$ die Menge der Punkte von $\mathfrak{R}$ in dieser Ebene, weiter in der Ebene $\{\mathfrak{F}\}$ mit $\{\mathfrak{L}\}$ die Stützgerade an $\mathfrak{F}$ mit der Normale $(\alpha', \beta', \gamma')$, mit $\mathfrak{L}$ die Menge der Punkte von $\mathfrak{F}$ in dieser Geraden. Der Punkt p ist dann endlich derjenige Punkt von $\mathfrak{L}$, von dem aus in der Richtung $(\alpha'', \beta'', \gamma'')$ kein weiterer Punkt von $\mathfrak{L}$ liegt.

In analoger Weise wie die Bereiche $\{\mathfrak{F}\}, \mathfrak{F}, \{\mathfrak{L}\}, \mathfrak{L}, p$ für $\mathfrak{R}$ mögen in bezug auf genau dieselbe orthogonale Substitution S für jeden Bezirk $\mathfrak{R}_i$ ($i = 1, 2, \dots, m$) die Bereiche $\{\mathfrak{F}_i\}, \mathfrak{F}_i, \{\mathfrak{L}_i\}, \mathfrak{L}_i, p_i$ definiert sein, so daß also insbesondere p_i der zu S gehörende extreme Punkt von $\mathfrak{R}_i$ wird.

Fragen wir nun, wie für den Punkt p eine Darstellung

$$(84) \quad p = s_1 q_1 + s_2 q_2 + \dots + s_m q_m$$

möglich ist, wobei q_i jedesmal ein Punkt aus $\mathfrak{R}_i$ sein soll.

Indem wir eine orthogonale Koordinatentransformation zulassen, können wir ohne Beeinträchtigung der Allgemeinheit annehmen, die Richtungen $(\alpha'', \beta'', \gamma'')$, $(\alpha', \beta', \gamma')$, (α, β, γ) seien die der positiven x, y, z -Achse. Die Ebene $\{\mathfrak{F}\}$, welche p enthält, ist dann durch $z = H(0, 0, 1)$

bestimmt; für einen Punkt $q_i(x_i, y_i, z_i)$ aus $\mathfrak{R}_i$ gilt jedenfalls $z_i \leq H_i(0, 0, 1)$. Nun haben wir insbesondere

$$(85) \quad H(0, 0, 1) = s_1 H_1(0, 0, 1) + s_2 H_2(0, 0, 1) + \cdots + s_m H_m(0, 0, 1).$$

Eine Relation, wie wir sie in (84) fordern, kann danach gewiß nur in der Weise statthaben, daß dabei für jeden Punkt q_i stets $z_i = H_i(0, 0, 1)$ ist, d. h. daß q_i in der Ebene $\{\mathfrak{F}_i\}$, als Punkt von $\mathfrak{R}_i$ also in $\mathfrak{F}_i$ liegt.

Nach § 12 besitzt die Funktion $H(u, v, w) - wH(0, 0, 1)$ und entsprechend besitzen die Ausdrücke $H_i(u, v, w) - wH_i(0, 0, 1)$ bei festgehaltenen Werten u, v für ein unbegrenzt wachsendes w jedesmal einen bestimmten Grenzwert; wir bezeichnen diese Grenzwerte mit $H'(u, v)$, $H'_i(u, v)$. Aus $H(u, v, w) = \sum s_i H_i(u, v, w)$ und der daraus entnommenen speziellen Gleichung (85) folgt dann $H'(u, v) = \sum s_i H'_i(u, v)$.

Die Gerade $\{\mathfrak{L}\}$ der Ebene $\{\mathfrak{F}\}$ wird durch $y = H'(0, 1)$ bestimmt und hat daher für p weiter dieser Wert von y statt. Für einen Punkt q_i aus $\mathfrak{F}_i$ gilt jedenfalls $y_i \leq H'_i(0, 1)$. Nun ist insbesondere

$$(86) \quad H'(0, 1) = s_1 H'_1(0, 1) + s_2 H'_2(0, 1) + \cdots + s_m H'_m(0, 1).$$

Danach kann eine Relation, wie wir sie in (84) fordern, jedenfalls nur so statthaben, daß dabei für q_i jedesmal $y_i = H'_i(0, 1)$ ist, d. h. daß q_i auf der Geraden $\{\mathfrak{L}_i\}$, also als Punkt von $\mathfrak{F}_i$ notwendig in $\mathfrak{L}_i$ liegt.

Weiter besitzen nach § 12 die Ausdrücke $H'(u, v) - vH'(0, 1)$, $H'_i(u, v) - vH'_i(0, 1)$ bei festem u für ein unbegrenzt wachsendes v bestimmte Grenzwerte, die wir $H''(u)$, $H''_i(u)$ nennen wollen. Aus $H'(u, v) = \sum s_i H'_i(u, v)$ und der daraus abgeleiteten speziellen Gleichung (86) folgt dann $H''(u) = \sum s_i H''_i(u)$.

Für den Punkt p in $\mathfrak{L}$ ist nun endlich $x = H''(1)$. Für einen Punkt q_i aus $\mathfrak{L}_i$ gilt jedenfalls $x_i \leq H''_i(1)$. Nun haben wir insbesondere

$$(87) \quad H''(1) = s_1 H''_1(1) + s_2 H''_2(1) + \cdots + s_m H''_m(1).$$

Danach kann endlich die Relation (84) nur noch so statthaben, daß dabei für q_i jedesmal $x_i = H''_i(1)$ ist, d. h. daß q_i mit dem zur Substitution S gehörenden extremen Punkte p_i von $\mathfrak{R}_i$ identisch wird. Mit den Punkten $q_i = p_i$ aber besteht nach den Gleichungen (85), (86), (87) in der Tat die Formel (84).

Auf diese Weise sind wir zu dem folgenden Resultate gekommen:

Ein beliebiger extremer Punkt p des Bezirks $\mathfrak{R}$ mit der Stützebenenfunktion $H = \sum s_i H_i$, wobei die Werte s_i sämtlich > 0 sind, läßt sich stets auf eine und nur eine Weise in die Form $p = \sum s_i p_i$ setzen, so daß dabei jedesmal p_i ein Punkt von $\mathfrak{R}_i$ ist. Gehört p als extremer Punkt von $\mathfrak{R}$ insbesondere zu einer orthogonalen Substitution S , so muß dabei p_i der zu derselben Substitution S gehörende extreme Punkt von $\mathfrak{R}_i$ werden.

Wir sahen bereits in § 17 ((76) und (77)), daß, wenn für zwei Punkte

p und p^* Darstellungen der hier verlangten Form (84) bestehen, daraus sogleich entsprechende Darstellungen für jeden Punkt der p und p^* verbindenden Strecke folgen. Von dem soeben abgeleiteten Satze aus gelangen wir, auf diese Bemerkung gestützt, nacheinander zu den weiteren Beziehungen:

$$(88) \quad \mathfrak{L} = \sum s_i \mathfrak{L}_i, \quad \mathfrak{F} = \sum s_i \mathfrak{F}_i, \quad \mathfrak{R} = \sum s_i \mathfrak{R}_i \quad (i = 1, 2, \dots, m),$$

deren Sinn sein soll, daß für jeden beliebigen Punkt q aus $\mathfrak{L}$, bzw. aus $\mathfrak{F}$, bzw. aus $\mathfrak{R}$ stets eine Darstellung $q = \sum s_i q_i$ möglich ist, so daß dabei q_i jedesmal einen Punkt aus $\mathfrak{L}_i$, bzw. aus $\mathfrak{F}_i$, bzw. aus $\mathfrak{R}_i$ bedeutet.

Sind die Grundbezirke $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ speziell m Punkte, oder sind sie ganz in m parallelen Geraden oder sind sie ganz in m parallelen Ebenen gelegen, so ist jeder Bezirk $\mathfrak{R}$ der aus ihnen entspringenden Schar ein Punkt, oder in einer zu den m Geraden parallelen Geraden oder in einer zu den m Ebenen parallelen Ebene gelegen.

§ 19. Polyederscharen.

Wir wollen jetzt die Annahme verfolgen, daß $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ speziell lauter konvexe Bezirke je mit einer *endlichen* Anzahl von extremen Punkten, also Polyeder, Polygone, Strecken, Punkte seien. Alsdann gilt der Satz:

Jeder beliebige Bezirk $\mathfrak{R} = \sum s_i \mathfrak{R}_i$ der Schar mit den Grundbezirken $\mathfrak{R}_i$ weist notwendig nur eine endliche Anzahl von extremen Punkten auf, kann also gleichfalls nur ein Polyeder, Polygon, eine Strecke oder ein Punkt sein.

Dieser Satz ist unmittelbar daraus einleuchtend, daß jeder extreme Punkt von $\mathfrak{R}$ nach § 18 in der Form $\sum s_i p_i$ auftreten muß, wobei p_i jedesmal einen extremen Punkt von $\mathfrak{R}_i$ vorstellt. Wir können weiter feststellen:

Alle inwendigen Bezirke der Schar besitzen untereinander die gleiche Anzahl von Ecken mit den gleichen zugehörigen Normalensektoren in der Kugel $\mathfrak{G} (x^2 + y^2 + z^2 \leq 1)$.

In der Tat, setzen wir jetzt in $\mathfrak{R} = \sum s_i \mathfrak{R}_i$ die sämtlichen m Parameter $s_i > 0$ voraus. Es sei p eine beliebige Ecke, also ein extremer Punkt von $\mathfrak{R}$, so ist die Relation $p = \sum s_i p_i$ mit Punkten p_i aus den Bezirken $\mathfrak{R}_i$ nach § 18 auf eine einzige Weise herzustellen, wobei dann jedesmal p_i wieder einen extremen Punkt, also unter den gegenwärtigen Umständen eine gewisse Ecke von $\mathfrak{R}_i$ bedeutet. Ist dann (α, β, γ) Normale einer Stützebene durch p an $\mathfrak{R}$, so muß die Stützebene an $\mathfrak{R}_i$ mit dieser Normale (α, β, γ) jedesmal durch den Punkt p_i gehen. Bezeichnen wir mit $N(p)$ den (in der Kugel $\mathfrak{G}$ konstruierten) Normalensektor der

Ecke p von $\mathfrak{R}$, mit $N_i(p_i)$ den Normalensektor der Ecke p_i von $\mathfrak{R}_i$, so muß danach der ganze Sektor $N(p)$ stets in jedem einzelnen Sektor $N_i(p_i)$ enthalten sein, es müssen somit die hier miteinander in Verbindung tretenden Ecken p_1 von $\mathfrak{R}_1$, p_2 von $\mathfrak{R}_2$, ..., p_m von $\mathfrak{R}_m$ jedenfalls derart beschaffen sein, daß ihre zugehörigen Normalensektoren $N_1(p_1)$, $N_2(p_2)$, ..., $N_m(p_m)$ gewisse innere Punkte alle gemeinsam haben.

Umgekehrt bedeute p_i für $i = 1, 2, \dots, m$ jedesmal eine Ecke von $\mathfrak{R}_i$ und es mögen dabei die m zugehörigen Normalensektoren $N_i(p_i)$ sämtlich gewisse innere Punkte gemeinsam haben. Alsdann ist das ihnen allen gemeinsame Gebiet (s. den Schluß von § 11) wieder ein gewisser konvexer Körper N , der ebenfalls aus lauter Radien der Kugel $\mathfrak{G}$ von o aus bestehen wird. Ist (α, β, γ) die Richtung irgendeines solchen Radius von N , der von o aus ins Innere von N und also damit zugleich ins Innere eines jeden der Sektoren $N_i(p_i)$ eintritt, so ist die Stützebene an $\mathfrak{R}_i$ mit der Normale (α, β, γ) jedesmal eine Eckstützebene durch p_i an $\mathfrak{R}_i$, enthält also von $\mathfrak{R}_i$ allein den Punkt p_i . Die Stützebene mit der Normale (α, β, γ) an $\mathfrak{R}$ kann alsdann von $\mathfrak{R}$ allein den Punkt $p = \sum s_i p_i$ enthalten; also ist dieser Punkt p notwendig ein extremer Punkt, somit eine Ecke von $\mathfrak{R}$ und N stellt den Normalensektor dieser Ecke p von $\mathfrak{R}$ vor.

Wir bemerken nun noch, daß nach dem Satze am Schlusse von § 15 die Normalensektoren zu den sämtlichen Ecken des Bezirks $\mathfrak{R}$ oder eines der Bezirke $\mathfrak{R}_i$ jedesmal die Kugel $\mathfrak{G}$ vollständig erfüllen müssen, und können alsdann die Normalensektoren der sämtlichen Ecken von $\mathfrak{R}$ offenbar in folgender Weise aus den Grundbezirken $\mathfrak{R}_i$ ermitteln. Es möge p_i nacheinander die sämtlichen Ecken von $\mathfrak{R}_i$ durchlaufen. Die zugehörigen Normalensektoren $N_i(p_i)$ erfüllen jedesmal zusammengenommen die ganze Kugel $\mathfrak{G}$, und erscheint danach die Kugel $\mathfrak{G}$, den Werten $i = 1, 2, \dots, m$ entsprechend, auf m Weisen in Sektoren zerlegt. Jeder solche Radius von $\mathfrak{G}$, welcher bei keiner einzigen dieser m Zerlegungen als begrenzender Radius eines Sektors $N_i(p_i)$ auftritt, ist dann ein innerer Radius eines ganz bestimmten Sektors $N_1(p_1)$, eines ganz bestimmten Sektors $N_2(p_2)$, ..., eines ganz bestimmten Sektors $N_m(p_m)$, so daß die betreffenden m Sektoren dann notwendig gewisse innere Punkte gemein haben. Danach erhalten wir genau die Zerschneidung der Kugel $\mathfrak{G}$ in die gesuchten Normalensektoren $N(p)$ der sämtlichen Ecken p des Bezirks $\mathfrak{R}$, indem wir alle Begrenzungen, welche bei jenen m einzelnen Zerlegungen von $\mathfrak{G}$ auftreten, gleichzeitig zu einer einzigen Zerlegung der Kugel $\mathfrak{G}$ verwenden. Diese Sektoren $N(p)$ fallen auf diese Weise in der Tat identisch aus für alle inwendigen Bezirke $\mathfrak{R}$ der Schar.

Wenn $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ nicht m ebene Bezirke (Polygone, Strecken,

Punkte) in m parallelen Ebenen sind, ist jeder inwendige Bezirk $\mathfrak{R}$ der aus ihnen entspringenden Schar ein Polyeder und wollen wir von der Schar als einer *Polyederschar* sprechen. Die Randbezirke der Schar, speziell die Grundbezirke selbst, können dabei auch Polygone, Strecken, Punkte sein; sie lassen sich dann in der Regel als Grenzfälle der inwendigen Bezirke auffassen.

Wir setzen jetzt die Schar der Bezirke $\mathfrak{R} = \sum s_i \mathfrak{R}_i$ ($s_i \geq 0$, $i = 1, 2, \dots, m$) als Polyederschar voraus. Aus dem zuletzt bewiesenen Satze können wir dann weiter ableiten, daß die inwendigen Polyeder der Schar auch in bezug auf die Normalen und Anordnungen der Seitenflächen und Kanten sämtlich übereinstimmen.

In der Tat, betrachten wir einen beliebigen inwendigen Bezirk $\mathfrak{R}$ der Schar. Es sei $\mathfrak{R}$ ein Polyeder mit ν Seitenflächen, die wir in irgendeiner Weise numeriert $\mathfrak{F}^{(1)}, \mathfrak{F}^{(2)}, \dots, \mathfrak{F}^{(\nu)}$ nennen. Es sei $(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)})$ die Normale von $\mathfrak{F}^{(\tau)}$ für $\tau = 1, 2, \dots, \nu$. Sodann sei $\mathfrak{F}^{(\tau)}$ ein Polygon mit μ_τ Seiten, die wir in irgendeiner Weise numeriert mit $\mathfrak{L}^{(\tau 1)}, \mathfrak{L}^{(\tau 2)}, \dots, \mathfrak{L}^{(\tau \mu_\tau)}$ bezeichnen. Es sei $(\alpha^{(\tau \sigma)}, \beta^{(\tau \sigma)}, \gamma^{(\tau \sigma)})$ die Normale der Seite $\mathfrak{L}^{(\tau \sigma)}$ in $\mathfrak{F}^{(\tau)}$ für $\sigma = 1, 2, \dots, \mu_\tau$. Endlich bezeichnen wir die Endpunkte der Strecke $\mathfrak{L}^{(\tau \sigma)}$ mit $p^{(\tau \sigma 1)}$ und $p^{(\tau \sigma 2)}$, und es sei $(\alpha^{(\tau \sigma \varrho)}, \beta^{(\tau \sigma \varrho)}, \gamma^{(\tau \sigma \varrho)})$ für $\varrho = 1, 2$ diejenige Richtung in dieser Strecke, welche von $p^{(\tau \sigma \varrho)}$ aus ihr herausführt. Wir setzen ferner jedesmal

$$(89) \quad \begin{aligned} \xi^{(\tau \sigma \varrho)} &= \alpha^{(\tau \sigma \varrho)} x + \beta^{(\tau \sigma \varrho)} y + \gamma^{(\tau \sigma \varrho)} z, \\ \eta^{(\tau \sigma)} &= \alpha^{(\tau \sigma)} x + \beta^{(\tau \sigma)} y + \gamma^{(\tau \sigma)} z, \\ \xi^{(\tau)} &= \alpha^{(\tau)} x + \beta^{(\tau)} y + \gamma^{(\tau)} z, \end{aligned}$$

und bezeichnen die durch diese drei Gleichungen bestimmte orthogonale Substitution mit $S^{(\tau \sigma \varrho)}$.

Die Punkte $p^{(\tau \sigma \varrho)}$ werden die Ecken von $\mathfrak{R}$. Eine jede Ecke $p^{(\tau \sigma \varrho)}$ gehört in $\mathfrak{F}^{(\tau)}$ außer der Kante $\mathfrak{L}^{(\tau \sigma)}$ stets noch einer zweiten Kante $\mathfrak{L}^{(\tau \sigma')}$

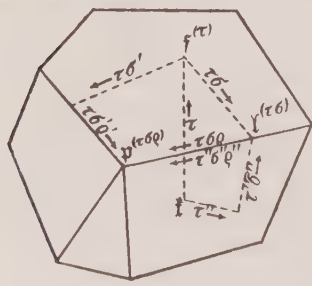


Fig. 2.

zu und gibt es also unter jenen Indexsystemen zu dem Systeme τ, σ, ϱ stets ein bestimmtes zweites System τ, σ', ϱ' ($\sigma' \neq \sigma$), wobei die Ecke $p^{(\tau \sigma' \varrho')} = p^{(\tau \sigma \varrho)}$ ist. Die Beziehung der zwei Systeme τ, σ, ϱ und τ, σ', ϱ' zueinander ist offenbar eine gegenseitige. Die zugehörigen orthogonalen Substitutionen $S^{(\tau \sigma \varrho)}$ und $S^{(\tau \sigma' \varrho')}$ stimmen in ihrer Form $\xi^{(\tau)}$ überein, während in der Ebene $\xi^{(\tau)} = 0$ das System der $\xi^{(\tau \sigma \varrho)}$ - und der $\eta^{(\tau \sigma)}$ -Achse einerseits und das System

der $\xi^{(\tau \sigma' \varrho')}$ - und der $\eta^{(\tau \sigma')}$ -Achse andererseits nicht kongruent sind, so daß die Determinanten dieser zwei Substitutionen stets entgegengesetzt gleich sind.

Ferner wird eine jede Kante $\mathfrak{L}^{(\tau\sigma)}$ von $\mathfrak{F}^{(\tau)}$ aus der Ebene von $\mathfrak{F}^{(\tau)}$ stets durch die Ebene einer bestimmten anderen Seitenfläche $\mathfrak{F}^{(\tau')}$ von $\mathfrak{R}$ herausgeschnitten und gehört dann auch dieser Seitenfläche als Kante zu; danach gibt es unter jenen Indexsystemen zu jedem Systeme τ, σ, ϱ ein bestimmtes anderes System $\tau'', \sigma'', \varrho''$ ($\tau'' \neq \tau$), wobei $\mathfrak{L}^{(\tau''\sigma'')} = \mathfrak{L}^{(\tau\sigma)}$, $\mathfrak{p}^{(\tau''\sigma''\varrho'')} = \mathfrak{p}^{(\tau\sigma\varrho)}$ und daher auch $\xi^{(\tau''\sigma''\varrho'')} = \xi^{(\tau\sigma\varrho)}$ ist. Die Beziehung solcher zwei Systeme τ, σ, ϱ und $\tau'', \sigma'', \varrho''$ zueinander ist eine gegenseitige. Die zugehörigen Substitutionen $S^{(\tau\sigma\varrho)}$ und $S^{(\tau''\sigma''\varrho'')}$ stimmen in ihrer ξ -Form überein, dagegen sind in der gemeinsamen Ebene $\xi^{(\tau\sigma\varrho)} = \xi^{(\tau''\sigma''\varrho'')} = 0$ das System der $\eta^{(\tau\sigma)}$ - und der $\xi^{(\tau)}$ -Achse einerseits und das System der $\eta^{(\tau''\sigma'')}$ - und der $\xi^{(\tau')}$ -Achse andererseits nicht kongruent, so daß die Determinanten dieser zwei Substitutionen stets entgegengesetzt gleich sind.

Nehmen wir jetzt ein beliebiges anderes inwendiges Polyeder $\mathfrak{R}^*$ der Schar, so muß es in $\mathfrak{R}^*$ nach dem vorhin bewiesenen Satze eine Ecke $\mathfrak{p}^*$ geben mit dem gleichen Normalensektor in der Kugel $\mathfrak{G}$ wie die Ecke $\mathfrak{p}^{(\tau\sigma\varrho)}$ von $\mathfrak{R}$. Alsdann muß der Projektionsraum des Polyeders $\mathfrak{R}$ von $\mathfrak{p}^{(\tau\sigma\varrho)}$ aus durch die Translation von $\mathfrak{p}^{(\tau\sigma\varrho)}$ nach $\mathfrak{p}^*$ unmittelbar in den Projektionsraum des Polyeders $\mathfrak{R}^*$ von $\mathfrak{p}^*$ aus übergehen. Danach muß $\mathfrak{R}^*$ eine Seitenfläche $\mathfrak{F}^*$ mit der Normale $(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)})$ und in dieser eine Kante $\mathfrak{L}^*$ mit der Normale $(\alpha^{(\tau\sigma)}, \beta^{(\tau\sigma)}, \gamma^{(\tau\sigma)})$ und dabei endlich $\mathfrak{p}^*$ als einen solchen Endpunkt von $\mathfrak{L}^*$ besitzen, daß in der Richtung $(\alpha^{(\tau\sigma\varrho)}, \beta^{(\tau\sigma\varrho)}, \gamma^{(\tau\sigma\varrho)})$ von $\mathfrak{p}^*$ kein weiterer Punkt von $\mathfrak{L}^*$ liegt; ferner muß in $\mathfrak{p}^*$ eine zweite Kante von $\mathfrak{F}^*$ mit der Normale $(\alpha^{(\tau\sigma')}, \beta^{(\tau\sigma')}, \gamma^{(\tau\sigma')})$ enden und muß $\mathfrak{L}^*$ einer zweiten Seitenfläche von $\mathfrak{R}^*$ mit der Normale $(\alpha^{(\tau')}, \beta^{(\tau')}, \gamma^{(\tau')})$ zugehören. Da in derselben Weise die bei $\mathfrak{R}^*$ vorkommenden Normalen der Seitenflächen und Kanten sich auch sämtlich bei $\mathfrak{R}$ vorfinden müssen, erkennen wir, daß die mit Hilfe irgendeines inwendigen Polyeders $\mathfrak{R}$ der Schar eingeführten Zahlen $\nu, \mu_1, \mu_2, \dots, \mu_\nu$, Richtungen

$$\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)}; \quad \alpha^{(\tau\sigma)}, \beta^{(\tau\sigma)}, \gamma^{(\tau\sigma)}; \quad \alpha^{(\tau\sigma\varrho)}, \beta^{(\tau\sigma\varrho)}, \gamma^{(\tau\sigma\varrho)} \\ (\tau = 1, 2, \dots, \nu; \sigma = 1, 2, \dots, \mu_\tau; \varrho = 1, 2)$$

und Zuordnungen $\tau, \sigma, \varrho; \tau, \sigma', \varrho'; \tau'', \sigma'', \varrho''$ eine unveränderliche Bedeutung für alle inwendigen Polyeder der Schar besitzen.

Wir wollen noch zusehen, wie die hier genannten Größen aus den Grundbezirken $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ selbst hergeleitet werden können. Wir verstehen für jeden Bezirk $\mathfrak{R}_i$ ($i = 1, 2, \dots, m$) unter $\{\mathfrak{F}_i^{(\tau)}\}$ die Stützebene an $\mathfrak{R}_i$ mit der Normale $(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)})$, unter $\mathfrak{F}_i^{(\tau)}$ die Menge der Punkte von $\mathfrak{R}_i$ in $\{\mathfrak{F}_i^{(\tau)}\}$, sodann in dieser Ebene $\{\mathfrak{F}_i^{(\tau)}\}$ unter $\{\mathfrak{L}_i^{(\tau\sigma)}\}$ die Stützgerade an $\mathfrak{F}_i^{(\tau)}$ mit der Normale $(\alpha^{(\tau\sigma)}, \beta^{(\tau\sigma)}, \gamma^{(\tau\sigma)})$, unter $\mathfrak{L}_i^{(\tau\sigma)}$ die Menge der Punkte von $\mathfrak{F}_i^{(\tau)}$ in $\{\mathfrak{L}_i^{(\tau\sigma)}\}$, endlich unter $\mathfrak{p}_i^{(\tau\sigma\varrho)}$ denjenigen Punkt in $\mathfrak{L}_i^{(\tau\sigma)}$, von dem aus in der Richtung $(\alpha^{(\tau\sigma\varrho)}, \beta^{(\tau\sigma\varrho)}, \gamma^{(\tau\sigma\varrho)})$ kein weiterer

Punkt von $\mathfrak{L}_i^{(\tau\sigma)}$ liegt. Alsdann gelten für einen beliebigen inwendigen Bezirk $\mathfrak{R} = \sum s_i \mathfrak{R}_i$ der Schar nach § 18 stets die Beziehungen:

$$(90) \quad \mathfrak{F}^{(\tau)} = \sum s_i \mathfrak{F}_i^{(\tau)}, \quad \mathfrak{L}^{(\tau\sigma)} = \sum s_i \mathfrak{L}_i^{(\tau\sigma)}, \quad p^{(\tau\sigma\varrho)} = \sum s_i p_i^{(\tau\sigma\varrho)}.$$

Nun soll für $\mathfrak{R}$ jeder Bereich $\mathfrak{F}^{(\tau)}$ ($\tau = 1, 2, \dots, \nu$) wirklich ein Polygon sein und können nach der ersten Relation hier daher nicht $\mathfrak{F}_1^{(\tau)}$, $\mathfrak{F}_2^{(\tau)}, \dots, \mathfrak{F}_m^{(\tau)}$ lauter Strecken oder Punkte in m parallelen Geraden vorstellen, wir müssen also unter diesen m Bereichen entweder wenigstens ein Polygon oder wenigstens zwei einander nicht parallele Strecken vorfinden. Die sämtlichen Richtungen $(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)})$, die als äußere Normalen von Seitenflächen in $\mathfrak{R}$ auftreten, entstehen danach auf zweierlei Arten: *Erstens* haben wir darunter jede Richtung, welche bei irgendeinem der Grundbezirke $\mathfrak{R}_i$ als äußere Normale einer Seitenfläche vorkommt. So oft wir *zweitens* bei irgend zweien der Grundbezirke $\mathfrak{R}_i, \mathfrak{R}_j$ ($i \neq j$) zwei nicht parallele Kanten $\mathfrak{L}_i, \mathfrak{L}_j$ vorfinden, derart, daß die Ebene durch $\mathfrak{L}_i$ parallel zu $\mathfrak{L}_j$ Stützebene an $\mathfrak{R}_i$ und die Ebene durch $\mathfrak{L}_j$ parallel zu $\mathfrak{L}_i$ Stützebene an $\mathfrak{R}_j$ ist, und zwar diese beiden Stützebenen als solche mit gleicher (nicht mit entgegengesetzter) Normale auftreten, kommt die betreffende Normale stets ebenfalls unter den Richtungen $(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)})$ vor.

Hierbei haben wir in bezug auf diejenigen Bezirke $\mathfrak{R}_i$, welche Polygone oder Strecken bedeuten, noch folgendes zu beachten. Bei einem Polygon ist die Ebene des Polygons gleichzeitig Stützebene an dasselbe mit zwei einander entgegengesetzten Normalen und haben wir den Bereich des Polygons wie zwei Seitenflächen für das Polygon mit den betreffenden zwei Normalen aufzufassen, ferner die Seiten des Polygons als Kanten dieser Seitenflächen in Betracht zu ziehen. Bei einem Bezirk, der eine Strecke bedeutet, haben wir eine Kante eben in der Strecke selbst in Betracht zu ziehen.

Jeder Bereich $\mathfrak{L}^{(\tau\sigma)}$ für das Polyeder $\mathfrak{R}$ soll wirklich eine Strecke sein, und muß wegen der zweiten Gleichung in (90) dann unter den m Bereichen $\mathfrak{L}_1^{(\tau\sigma)}, \mathfrak{L}_2^{(\tau\sigma)}, \dots, \mathfrak{L}_m^{(\tau\sigma)}$ jedesmal wenigstens eine Strecke vorhanden sein. Danach haben wir unter den Richtungen $(\alpha^{(\tau\sigma)}, \beta^{(\tau\sigma)}, \gamma^{(\tau\sigma)})$ zu einem bestimmten Index τ eine jede solche, in $\xi^{(\tau)} = 0$ vorkommende Richtung, welche als äußere Normale einer Kante bei einem Bereiche $\mathfrak{F}_i^{(\tau)}$ in einer Ebene $\{\mathfrak{F}_i^{(\tau)}\}$ auftritt, und nur diese Richtungen.

§ 20. Volumen und Schwerpunkt in einer Polyederschar.

Betrachten wir wieder eine Polyederschar mit m Grundbezirken $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$. Wir suchen jetzt das Volumen und ferner die Koordinaten des Schwerpunkts für einen beliebigen Bezirk $\mathfrak{R} = \sum s_i \mathfrak{R}_i$ der Schar als Funktion der Parameter s_i darzustellen. Zu dem Ende werden wir

zunächst zeigen, wie ein solcher Bereich $\mathfrak{R}$ sich in gewisser Weise aus lauter Tetraedern aufbaut.

Wir halten an allen im vorigen Paragraphen eingeführten Bezeichnungen fest, und wir wollen uns zunächst die s_i sämtlich > 0 , also $\mathfrak{R}$ als inwendigen Bezirk der Schar denken. Wir nehmen einen Punkt $\mathfrak{f}$ im Raume beliebig an, es seien a, b, c die Koordinaten von $\mathfrak{f}$ und wir fällen von $\mathfrak{f}$ Lote auf die Ebenen $\{\mathfrak{F}^{(\tau)}\}$ der einzelnen Seitenflächen $\mathfrak{F}^{(\tau)}$ von $\mathfrak{R}$ (für $\tau = 1, 2, \dots, \nu$). Wir bezeichnen mit $C^{(\tau)}$ die Länge des Lotes von $\mathfrak{f}$ auf die Ebene $\{\mathfrak{F}^{(\tau)}\}$ und zwar noch mit negativem Vorzeichen versehen, falls diese Ebene den Punkt $\mathfrak{f}$ und das Innere von $\mathfrak{R}$ auf verschiedenen Seiten von sich liegen hat, so daß also in jedem Falle $C^{(\tau)}$ das Maximum von $\alpha^{(\tau)}(x - a) + \beta^{(\tau)}(y - b) + \gamma^{(\tau)}(z - c)$ für die Punkte x, y, z in $\mathfrak{R}$ bedeutet. Wir wollen bei einer beliebigen reellen Größe C unter $\text{sgn } C$ das Vorzeichen ± 1 von C verstehen, wenn $C \neq 0$ ist, oder den Wert Null, wenn $C = 0$ ist. Das ganze Polyeder $\mathfrak{R}$ setzt sich nun in gewisser Weise mit Hilfe der einzelnen Pyramiden $\mathfrak{f}\mathfrak{F}^{(\tau)}$ zusammen, welche $\mathfrak{f}$ als Spitze und die einzelnen Flächen $\mathfrak{F}^{(\tau)}$ als Grundflächen haben. Liegt $\mathfrak{f}$ im Inneren von $\mathfrak{R}$, so sind alle Größen $C^{(\tau)} > 0$ und überdecken jene Pyramiden zusammen genau den Bereich von $\mathfrak{R}$, wobei sie untereinander nur in den Begrenzungen gemeinsame Punkte haben. Entsprechendes gilt, wenn $\mathfrak{f}$ auf der Begrenzung von $\mathfrak{R}$ liegt, nur daß alsdann eine oder mehrere Größen $C^{(\tau)}$ gleich Null sind. Liegt $\mathfrak{f}$ außerhalb $\mathfrak{R}$, so überdecken diejenigen jener Pyramiden $\mathfrak{f}\mathfrak{F}^{(\tau)}$, welche Werten $C^{(\tau)} > 0$ entsprechen, das kleinste, $\mathfrak{f}$ und $\mathfrak{R}$ zugleich enthaltende Polyeder $(\mathfrak{f}, \mathfrak{R})$ (s. § 16), während sie untereinander nur in den Begrenzungen gemeinsame Punkte haben, und andererseits bilden diejenigen Pyramiden $\mathfrak{f}\mathfrak{F}^{(\tau)}$, welche Werten $C^{(\tau)} < 0$ entsprechen, (wenn wir die Punkte aus den Flächen $\mathfrak{F}^{(\tau)}$ selbst fortlassen), genau die *Kappe* von $\mathfrak{f}$ an $\mathfrak{R}$, wobei diese Pyramiden ebenfalls nur in den Begrenzungen zusammenstoßen. So oft endlich eine Größe $C^{(\tau)} = 0$ ist, reduziert sich die betreffende Pyramide $\mathfrak{f}\mathfrak{F}^{(\tau)}$ auf eine Polygonfläche, erlangt also ein Volumen $= 0$.

Wir können hiernach in jedem Falle für den Bereich des Polyeders $\mathfrak{R}$ eine Relation

$$(91) \quad [\mathfrak{R}] = \sum \text{sgn } C^{(\tau)} \cdot [\mathfrak{f}\mathfrak{F}^{(\tau)}] \quad (\tau = 1, 2, \dots, \nu)$$

in folgendem Sinne behaupten: Bringen wir jeden Punkt des Raumes, der irgendwie in die Bereiche $\mathfrak{f}\mathfrak{F}^{(\tau)}$ zu liegen kommt, bei jedem Bereich $\mathfrak{f}\mathfrak{F}^{(\tau)}$, dem er angehört, mit dem Gewicht $\text{sgn } C^{(\tau)}$ in Rechnung, so resultiert hieraus (von den Punkten in einer endlichen Anzahl von Polygonflächen abgesehen) für jeden Punkt in $\mathfrak{R}$ ein Gesamtgewicht $= 1$, für jeden Punkt außerhalb $\mathfrak{R}$ ein Gesamtgewicht $= 0$. Betrachten wir nun irgendeine Funktion Φ der Koordinaten x, y, z und bezeichnen den Wert des Raum-

integrals $\iiint \Phi dx dy dz$ über einen Raum $\mathfrak{R}$, (wenn diesem Integrale eine Bedeutung zukommt), mit $J(\mathfrak{R})$, so besteht alsdann zwischen den Werten $J(\mathfrak{R})$ und $J(\mathfrak{f}\mathfrak{F}^{(\tau)})$ für das Polyeder $\mathfrak{R}$ und für die einzelnen Pyramiden $\mathfrak{f}\mathfrak{F}^{(\tau)}$ notwendig der folgende Zusammenhang:

$$(92) \quad J(\mathfrak{R}) = \sum \text{sgn } C^{(\tau)} \cdot J(\mathfrak{f}\mathfrak{F}^{(\tau)}) \quad (\tau = 1, 2, \dots, \nu).$$

Wir nehmen weiter in einer jeden Ebene $\{\mathfrak{F}^{(\tau)}\}$ einen Punkt $\mathfrak{f}^{(\tau)}$ beliebig an und fällen von ihm die Lote auf alle Geraden $\{\mathfrak{L}^{(\tau\sigma)}\}$ der Kanten $\mathfrak{L}^{(\tau\sigma)}$ von $\mathfrak{F}^{(\tau)}$ (für $\sigma = 1, 2, \dots, \mu_\tau$). Es sei $B^{(\tau\sigma)}$ die Länge des Lotes von $\mathfrak{f}^{(\tau)}$ auf $\{\mathfrak{L}^{(\tau\sigma)}\}$ und zwar noch mit negativem Vorzeichen versehen, wenn diese Gerade den Punkt $\mathfrak{f}^{(\tau)}$ und das Inwendige von $\mathfrak{F}^{(\tau)}$ auf verschiedenen Seiten von sich liegen hat. Als dann setzt sich das Polygon $\mathfrak{F}^{(\tau)}$ aus den einzelnen Dreiecken $\mathfrak{f}^{(\tau)}\mathfrak{L}^{(\tau\sigma)}$ mit $\mathfrak{f}^{(\tau)}$ als Spitze und den Kanten $\mathfrak{L}^{(\tau\sigma)}$ als Grundlinien dergestalt additiv und subtraktiv zusammen, daß wir für die Bereiche dieser ebenen Flächen die Relation

$$(93) \quad [\mathfrak{F}^{(\tau)}] = \sum \text{sgn } B^{(\tau\sigma)} \cdot [\mathfrak{f}^{(\tau)}\mathfrak{L}^{(\tau\sigma)}] \quad (\sigma = 1, 2, \dots, \mu_\tau)$$

in einem ganz entsprechenden Sinne hinschreiben können, wie vorhin die Formel (91) aufgestellt wurde. Dieser Zusammensetzung des Polygons $\mathfrak{F}^{(\tau)}$ entspricht dann ein gewisser Aufbau der Pyramide $\mathfrak{f}\mathfrak{F}^{(\tau)}$ aus den einzelnen Tetraedern $\mathfrak{f}\mathfrak{f}^{(\tau)}\mathfrak{L}^{(\tau\sigma)}$ mit $\mathfrak{f}$ als Spitze und den Dreiecken $\mathfrak{f}^{(\tau)}\mathfrak{L}^{(\tau\sigma)}$ als Grundflächen, und können wir dadurch die Formel (92) weiter zu der Gleichung

$$(94) \quad J(\mathfrak{R}) = \sum \text{sgn } (B^{(\tau\sigma)} C^{(\tau)}) \cdot J(\mathfrak{f}\mathfrak{f}^{(\tau)}\mathfrak{L}^{(\tau\sigma)}) \quad \left(\begin{array}{l} \tau = 1, 2, \dots, \nu \\ \sigma = 1, 2, \dots, \mu_\tau \end{array} \right)$$

entwickeln.

Endlich nehmen wir noch in jeder Geraden $\{\mathfrak{L}^{(\tau\sigma)}\}$ einen Punkt $\mathfrak{l}^{(\tau\sigma)}$ beliebig an; dabei treffen wir die Bestimmung, daß bei je zwei Indexsystemen $\tau\sigma$ und $\tau''\sigma''$, wofür gemäß § 19 sich $\mathfrak{L}^{(\tau''\sigma'')} = \mathfrak{L}^{(\tau\sigma)}$ erweist, stets auch $\mathfrak{l}^{(\tau''\sigma'')} = \mathfrak{l}^{(\tau\sigma)}$ identisch gewählt werde. Es sei sodann $A^{(\tau\sigma\varrho)}$ die Länge der Entfernung von $\mathfrak{l}^{(\tau\sigma)}$ nach der Ecke $\mathfrak{p}^{(\tau\sigma\varrho)}$ von $\mathfrak{L}^{(\tau\sigma)}$ (für $\varrho = 1, 2$) und zwar noch mit negativem Vorzeichen versehen, wenn $\mathfrak{l}^{(\tau\sigma)}$ auf entgegengesetzter Seite der Ecke $\mathfrak{p}^{(\tau\sigma\varrho)}$ liegt als der zweite Endpunkt der Strecke $\mathfrak{L}^{(\tau\sigma)}$. Die Strecke $\mathfrak{L}^{(\tau\sigma)}$ setzt sich aus den zwei Stücken $\mathfrak{l}^{(\tau\sigma)}\mathfrak{p}^{(\tau\sigma\varrho)}$ ($\varrho = 1, 2$) dergestalt zusammen, daß wir für diese in der Geraden $\{\mathfrak{L}^{(\tau\sigma)}\}$ gelegenen Bereiche die Relation

$$(95) \quad [\mathfrak{L}^{(\tau\sigma)}] = \sum \text{sgn } A^{(\tau\sigma\varrho)} \cdot [\mathfrak{l}^{(\tau\sigma)}\mathfrak{p}^{(\tau\sigma\varrho)}] \quad (\varrho = 1, 2)$$

in einem entsprechenden Sinne haben, wie die Formel (93) einen Zusammenhang zwischen ebenen Bereichen darstellte. Aus dieser Zerlegung (95) folgt eine entsprechende Zerlegung des Dreiecks $\mathfrak{f}^{(\tau)}\mathfrak{L}^{(\tau\sigma)}$ in zwei Dreiecke $\mathfrak{f}^{(\tau)}\mathfrak{l}^{(\tau\sigma)}\mathfrak{p}^{(\tau\sigma\varrho)}$, und resultiert daraus weiter ein Aufbau von $\mathfrak{F}^{(\tau)}$ aus den Dreiecken $\mathfrak{f}^{(\tau)}\mathfrak{l}^{(\tau\sigma)}\mathfrak{p}^{(\tau\sigma\varrho)}$:

$$(96) \quad [\mathfrak{F}^{(\tau)}] = \sum \operatorname{sgn} (A^{(\tau\sigma\varrho)} B^{(\tau\sigma)}) \cdot [f^{(\tau)} l^{(\tau\sigma)} p^{(\tau\sigma\varrho)}] \begin{pmatrix} \sigma = 1, 2, \dots, \mu_\tau \\ \varrho = 1, 2 \end{pmatrix},$$

und schließlich folgt ein Aufbau des Polyeders $\mathfrak{R}$ aus den Tetraedern $f^{(\tau)} l^{(\tau\sigma)} p^{(\tau\sigma\varrho)}$:

$$(97) \quad [\mathfrak{R}] = \sum \operatorname{sgn} (A^{(\tau\sigma\varrho)} B^{(\tau\sigma)} C^{(\tau)}) \cdot [f^{(\tau)} l^{(\tau\sigma)} p^{(\tau\sigma\varrho)}] \begin{pmatrix} \tau = 1, 2, \dots, \nu \\ \sigma = 1, 2, \dots, \mu_\tau \\ \varrho = 1, 2 \end{pmatrix}.$$

Diese letzte Zerlegung von $\mathfrak{R}$ liefert uns endlich für das Integral $J(\mathfrak{R})$ die Beziehung

$$(98) \quad J(\mathfrak{R}) = \sum \operatorname{sgn} (A^{(\tau\sigma\varrho)} B^{(\tau\sigma)} C^{(\tau)}) \cdot J(f^{(\tau)} l^{(\tau\sigma)} p^{(\tau\sigma\varrho)}).$$

Wir wollen das Volumen von $\mathfrak{R}$ mit V , den Flächeninhalt von $\mathfrak{F}^{(\tau)}$ mit $F^{(\tau)}$, die Länge von $\mathfrak{L}^{(\tau\sigma)}$ mit $L^{(\tau\sigma)}$ bezeichnen. Die Länge der Strecke $l^{(\tau\sigma)} p^{(\tau\sigma\varrho)}$ ist der Betrag von $A^{(\tau\sigma\varrho)}$, der Flächeninhalt des Dreiecks $f^{(\tau)} l^{(\tau\sigma)} p^{(\tau\sigma\varrho)}$ das Produkt dieser Länge in den Betrag von $\frac{1}{2} B^{(\tau\sigma)}$, das Volumen des Tetraeders $f^{(\tau)} l^{(\tau\sigma)} p^{(\tau\sigma\varrho)}$ das Produkt dieses Flächeninhalts in den Betrag von $\frac{1}{3} C^{(\tau)}$. Danach entsteht aus (98), wenn wir $\Phi = 1$ nehmen, insbesondere

$$(99) \quad V = \frac{1}{6} \sum_{\tau, \sigma, \varrho} A^{(\tau\sigma\varrho)} B^{(\tau\sigma)} C^{(\tau)} \begin{pmatrix} \tau = 1, 2, \dots, \nu \\ \sigma = 1, 2, \dots, \mu_\tau \\ \varrho = 1, 2 \end{pmatrix},$$

während wir aus (96) und (95) analog erhalten:

$$(100) \quad F^{(\tau)} = \frac{1}{2} \sum_{\sigma, \varrho} A^{(\tau\sigma\varrho)} B^{(\tau\sigma)}, \quad L^{(\tau\sigma)} = A^{(\tau\sigma 1)} + A^{(\tau\sigma 2)}.$$

Wir nehmen zweitens $\Phi = x$ an und wir wollen das Integral $\iiint x dx dy dz$ über das Polyeder $\mathfrak{R}$ durch $\mathfrak{X}$ bezeichnen. Nennen wir die x -Koordinate für $f, f^{(\tau)}, l^{(\tau\sigma)}, p^{(\tau\sigma\varrho)}$ bzw. $a, a^{(\tau)}, a^{(\tau\sigma)}, a^{(\tau\sigma\varrho)}$, so ist der Wert des entsprechenden Raumintegrals über den Bereich des Tetraeders $f^{(\tau)} l^{(\tau\sigma)} p^{(\tau\sigma\varrho)}$ gleich dem Produkt aus dem Volumen dieses Tetraeders in

$$\frac{1}{4} (a + a^{(\tau)} + a^{(\tau\sigma)} + a^{(\tau\sigma\varrho)}),$$

und geht daher aus (98) für $\Phi = x$ die Relation

$$(101) \quad \mathfrak{X} = \frac{1}{24} \sum_{\tau, \sigma, \varrho} (a + a^{(\tau)} + a^{(\tau\sigma)} + a^{(\tau\sigma\varrho)}) A^{(\tau\sigma\varrho)} B^{(\tau\sigma)} C^{(\tau)}$$

hervor.

Über die Hilfspunkte $f, f^{(\tau)}, l^{(\tau\sigma)}$ treffen wir nun die folgende Verfügung. Wir nehmen zunächst für jeden Index $i = 1, 2, \dots, m$ und alle in Betracht kommenden τ, σ einen Punkt f_i beliebig an, weiter einen Punkt $f_i^{(\tau)}$ beliebig in der Stützebene $\{\mathfrak{F}_i^{(\tau)}\}$ an $\mathfrak{R}_i$, einen Punkt $l_i^{(\tau\sigma)}$ beliebig in der Stützgeraden $\{\mathfrak{L}_i^{(\tau\sigma)}\}$ an $\mathfrak{F}_i^{(\tau)}$; für je zwei gemäß § 19 zugeordnete Systeme

τ, σ und τ'', σ'' , wobei $\{\mathfrak{L}_i^{(\tau''\sigma'')}\} = \{\mathfrak{L}_i^{(\tau\sigma)}\}$ ist, wählen wir aber stets $\mathfrak{L}_i^{(\tau''\sigma'')} = \mathfrak{L}_i^{(\tau\sigma)}$. Hernach führen wir für jeden beliebigen Bezirk $\mathfrak{R} = \sum s_i \mathfrak{R}_i$ der Schar die Punkte $\mathfrak{f}, \mathfrak{f}^{(\tau)}, \mathfrak{I}^{(\tau\sigma)}$ durch die Formeln

$$(102) \quad \mathfrak{f} = \sum s_i \mathfrak{f}_i, \mathfrak{f}^{(\tau)} = \sum s_i \mathfrak{f}_i^{(\tau)}, \mathfrak{I}^{(\tau\sigma)} = \sum s_i \mathfrak{I}_i^{(\tau\sigma)} \quad (i = 1, 2, \dots, m)$$

ein. Nach § 18 wird dann in der Tat, wie wir es fordern, jedesmal $\mathfrak{f}^{(\tau)}$ in der Stützebene $\{\mathfrak{F}^{(\tau)}\}$ an $\mathfrak{R}$ und $\mathfrak{I}^{(\tau\sigma)}$ in der Stützgeraden $\{\mathfrak{L}^{(\tau\sigma)}\}$ an $\mathfrak{F}^{(\tau)}$ liegen. Wir merken noch an, daß für die Ecken $\mathfrak{p}^{(\tau\sigma\varrho)}$ von $\mathfrak{R}$ die Formeln gelten:

$$(103) \quad \mathfrak{p}^{(\tau\sigma\varrho)} = \sum s_i \mathfrak{p}_i^{(\tau\sigma\varrho)} \quad (i = 1, 2, \dots, m).$$

Führen wir die Substitution $S^{(\tau\sigma\varrho)}$ gemäß (89) ein, so bedeutet $C^{(\tau)}$ die Differenz der Werte von $\xi^{(\tau)}$ für $\mathfrak{f}^{(\tau)}$ und $\mathfrak{f}$, sodann $B^{(\tau\sigma)}$ die Differenz der Werte von $\eta^{(\tau\sigma)}$ für $\mathfrak{I}^{(\tau\sigma)}$ und $\mathfrak{f}^{(\tau)}$, endlich $A^{(\tau\sigma\varrho)}$ die Differenz der Werte von $\xi^{(\tau\sigma\varrho)}$ für $\mathfrak{p}^{(\tau\sigma\varrho)}$ und $\mathfrak{I}^{(\tau\sigma)}$. Bezeichnen wir für $\mathfrak{R}_i$ die Differenz der $\xi^{(\tau)}$ -Werte von $\mathfrak{f}_i^{(\tau)}$ und $\mathfrak{f}_i$ mit $C_i^{(\tau)}$, der $\eta^{(\tau\sigma)}$ -Werte von $\mathfrak{I}_i^{(\tau\sigma)}$ und $\mathfrak{f}_i^{(\tau)}$ mit $B_i^{(\tau\sigma)}$, der $\xi^{(\tau\sigma\varrho)}$ -Werte von $\mathfrak{p}_i^{(\tau\sigma\varrho)}$ und $\mathfrak{I}_i^{(\tau\sigma)}$ mit $A_i^{(\tau\sigma\varrho)}$, so gelten zufolge (102) und (103) die Beziehungen:

$$(104) \quad C^{(\tau)} = \sum s_i C_i^{(\tau)}, B^{(\tau\sigma)} = \sum s_i B_i^{(\tau\sigma)}, A^{(\tau\sigma\varrho)} = \sum s_i A_i^{(\tau\sigma\varrho)}.$$

Auf Grund dieser Gleichungen erhalten wir aus (100):

$$(105) \quad L^{(\tau\sigma)} = \sum L_g^{(\tau\sigma)} s_g, F^{(\tau)} = \sum F_{gh}^{(\tau)} s_g s_h \quad (g, h = 1, 2, \dots, m),$$

wenn

$$(106) \quad L_g^{(\tau\sigma)} = A_g^{(\tau\sigma 1)} + A_g^{(\tau\sigma 2)}, F_{gh}^{(\tau)} = \frac{1}{2} \sum_{\sigma, \varrho} A_g^{(\tau\sigma\varrho)} B_h^{(\tau\sigma)}$$

gesetzt wird. Aus (99) erhalten wir

$$(107) \quad V = \sum V_{ghj} s_g s_h s_j \quad (g, h, j = 1, 2, \dots, m),$$

wenn wir

$$(108) \quad V_{ghj} = \frac{1}{6} \sum_{\tau, \sigma, \varrho} A_g^{(\tau\sigma\varrho)} B_h^{(\tau\sigma)} C_j^{(\tau)}$$

eingeführen. Das Volumen V des Polyeders $\mathfrak{R}$ stellt sich danach als eine homogene Funktion dritten Grades der Parameter $s_1, s_2, \dots, s_m$ dar.

Im speziellen können wir, wovon bisweilen Gebrauch zu machen sein wird, jeden Punkt $\mathfrak{f}_i$ im Bezirke $\mathfrak{R}_i$ selbst, jeden Punkt $\mathfrak{f}_i^{(\tau)}$ im Bezirke $\mathfrak{F}_i^{(\tau)}$ selbst, jeden Punkt $\mathfrak{I}_i^{(\tau\sigma)}$ im Bezirke $\mathfrak{L}_i^{(\tau\sigma)}$ selbst wählen. Bei Verwendung der Formeln (102) wird dann, wie aus den Relationen (88) in § 18 zu ersehen ist, bei beliebigen Werten $s_i \geq 0$ stets $\mathfrak{f}$ in $\mathfrak{R}$ selbst, $\mathfrak{f}^{(\tau)}$ in $\mathfrak{F}^{(\tau)}$ selbst, $\mathfrak{I}^{(\tau\sigma)}$ in $\mathfrak{L}^{(\tau\sigma)}$ selbst zu liegen kommen und werden infolgedessen die Größen $A^{(\tau\sigma\varrho)}, B^{(\tau\sigma)}, C^{(\tau)}$ stets sämtlich ≥ 0 ausfallen.

Bezeichnen wir weiter die x -Koordinate von $\mathfrak{f}_i, \mathfrak{f}_i^{(\tau)}, \mathfrak{I}_i^{(\tau\sigma)}, \mathfrak{p}_i^{(\tau\sigma\varrho)}$ mit $a_i, a_i^{(\tau)}, a_i^{(\tau\sigma)}, a_i^{(\tau\sigma\varrho)}$, so folgen aus (103) und (102) die Gleichungen:

$$(109) \quad a = \sum s_i a_i, a^{(\tau)} = \sum s_i a_i^{(\tau)}, a^{(\tau\sigma)} = \sum s_i a_i^{(\tau\sigma)}, a^{(\tau\sigma\varrho)} = \sum s_i a_i^{(\tau\sigma\varrho)},$$

und ersehen wir aus (101), daß das Integral $\mathfrak{K}$ für $\mathfrak{R}$ eine homogene Funktion vierten Grades in $s_1, s_2, \dots, s_m$ wird.

Die hier zugrunde gelegte Zerlegung eines Bezirks $\mathfrak{R}$ läßt sich offenbar auch auf die Randbezirke der Schar übertragen und bleiben dadurch alle hier entwickelten Formeln auch noch für solche Parameterwerte s_i gültig, die zum Teil oder sämtlich gleich Null sind.

Wir werden jetzt die hier aufgestellten Entwicklungen von Polyedern auf beliebige konvexe Körper auszudehnen suchen. Zu dem Ende müssen wir vor allem die Bildungen $F_{gh}^{(\tau)}$ und V_{ghj} einer eingehenderen Untersuchung unterwerfen.

§ 21. Das gemischte Volumen dreier Polyeder.

Die m vorausgesetzten Grundbezirke $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ brauchen nicht verschieden zu sein; wir können in dieser Reihe auch einen und denselben Bezirk mehrfach aufnehmen, wodurch nur die Bedeutung der Parameter $s_1, s_2, \dots, s_m$ etwas modifiziert wird, ohne daß die Schar $\sum s_i \mathfrak{R}_i$ in ihrer Gesamtheit sich verändert. Es wird deshalb genügen, die Bildungen $F_{12}^{(\tau)}$, V_{123} mit verschiedenen Indizes zu betrachten, wobei wir uns bei dem ersten Ausdrucke $m \geq 2$, bei dem zweiten $m \geq 3$ denken. Wir werden nun vor allem den Satz beweisen:

Der Wert des Ausdrucks $F_{12}^{(\tau)}$ ändert sich nicht, wenn man darin die Indizes 1 und 2 vertauscht; der Ausdruck V_{123} behält bei jeder Permutation der Indizes 1, 2, 3 seinen Wert bei.

Danach wird dann der Koeffizient von $s_1 s_2$ im Ausdrucke (105) von $F^{(\tau)}$ genau $= 2 F_{12}^{(\tau)}$, der Koeffizient von $s_1 s_2 s_3$ im Ausdrucke (107) von V genau $= 6 V_{123}$ und werden infolgedessen die Werte von $F_{12}^{(\tau)}$ und V_{123} ganz unabhängig von der Wahl der Hilfspunkte $\mathfrak{f}_i, \mathfrak{f}_i^{(\tau)}, \mathfrak{l}_i^{(\tau\sigma)}$ sein.

Um jenen Satz zu beweisen, bemerken wir zunächst die Möglichkeit, diese Hilfspunkte so einzurichten, daß die folgenden Umstände zutreffen:

I. Während über $\mathfrak{f}_i$ irgendwie verfügt sei, soll $\mathfrak{f}_i^{(\tau)}$ speziell den Fußpunkt des von $\mathfrak{f}_i$ auf die Ebene $\{\mathfrak{F}_i^{(\tau)}\}$ gefällten Lotes, ferner $\mathfrak{l}_i^{(\tau\sigma)}$ den Fußpunkt des von $\mathfrak{f}_i$ auf die Gerade $\{\mathfrak{L}_i^{(\tau\sigma)}\}$ gefällten Lotes (d. i. also zugleich den Fußpunkt des von $\mathfrak{f}_i^{(\tau)}$ auf $\{\mathfrak{L}_i^{(\tau\sigma)}\}$ gefällten Lotes) vorstellen. Bei dieser Festsetzung werden $A_i^{(\tau\sigma\varrho)}, B_i^{(\tau\sigma)}, C_i^{(\tau)}$ einfach die Differenzen der Koordinaten $\xi^{(\tau\sigma\varrho)}, \eta^{(\tau\sigma)}, \zeta^{(\tau)}$ für $\mathfrak{p}_i^{(\tau\sigma\varrho)}$ und $\mathfrak{f}_i$, und kommen bei Anwendung der Formeln (102) allgemein $\mathfrak{f}$ mit $\mathfrak{f}^{(\tau)}$, $\mathfrak{f}^{(\tau)}$ mit $\mathfrak{l}^{(\tau\sigma)}$, $\mathfrak{l}^{(\tau\sigma)}$ mit $\mathfrak{p}^{(\tau\sigma\varrho)}$ in drei Gerade zu liegen, die beziehlich parallel der $\zeta^{(\tau)}, \eta^{(\tau\sigma)}, \xi^{(\tau\sigma\varrho)}$ -Achse der Substitution $S^{(\tau\sigma\varrho)}$ sind, so daß z. B. für die x -Koordinaten der be-

treffenden Punkte die Gleichungen

(110) $\alpha^{(\tau)} - \alpha = \alpha^{(\tau)} C^{(\tau)}$, $\alpha^{(\tau\sigma)} - \alpha^{(\tau)} = \alpha^{(\tau\sigma)} B^{(\tau\sigma)}$, $\alpha^{(\tau\sigma\varrho)} - \alpha^{(\tau\sigma)} = \alpha^{(\tau\sigma\varrho)} A^{(\tau\sigma\varrho)}$ gelten werden.

II. Ferner sollen $\mathfrak{f}_1, \mathfrak{f}_2, \dots, \mathfrak{f}_m$ sämtlich in einen und denselben Punkt $\mathfrak{f}_0$ des Raumes fallen, dessen Koordinaten a_0, b_0, c_0 seien. Für den Punkt $\mathfrak{f}$ in bezug auf einen beliebigen Bezirk $\mathfrak{R}$ gilt dann $\mathfrak{f} = (s_1 + s_2 + \dots + s_m) \mathfrak{f}_0$.

Bei diesen beschränkenden Festsetzungen I. und II. würden schließlich allein die Koordinaten des Punktes $\mathfrak{f}_0$ willkürlich bleiben. Durch geeignete Verfügung über $\mathfrak{f}_0$ aber würden wir noch imstande sein, einem einzelnen Punkte $\mathfrak{f}_1^{(\tau\sigma)}$ eine beliebige Lage auf $\{\mathfrak{L}_1^{(\tau\sigma)}\}$ oder einem einzelnen Punkte $\mathfrak{f}_2^{(\tau)}$ eine beliebige Lage auf $\{\mathfrak{F}_2^{(\tau)}\}$ zu verschaffen oder $\mathfrak{f}$ beliebig anzunehmen.

Wir denken uns nun zunächst die Hilfspunkte $\mathfrak{f}_i, \mathfrak{f}_i^{(\tau)}, \mathfrak{f}_i^{(\tau\sigma)}$ noch in der ganzen Allgemeinheit wie in § 20. Die Größe

$$(111) \quad L_1^{(\tau\sigma)} = A_1^{(\tau\sigma 1)} + A_1^{(\tau\sigma 2)}$$

bedeutet die Länge der Kante $\mathfrak{L}_1^{(\tau\sigma)}$, ist also völlig unabhängig von der Lage von $\mathfrak{f}_1^{(\tau\sigma)}$, eine Tatsache, die auf den Gleichungen

$$(112) \quad \alpha^{(\tau\sigma 1)} + \alpha^{(\tau\sigma 2)} = 0, \quad \beta^{(\tau\sigma 1)} + \beta^{(\tau\sigma 2)} = 0, \quad \gamma^{(\tau\sigma 1)} + \gamma^{(\tau\sigma 2)} = 0$$

beruht. Wir können alsdann

$$(113) \quad F_{12}^{(\tau)} = \frac{1}{2} \sum_{\sigma, \varrho} A_1^{(\tau\sigma\varrho)} B_2^{(\tau\sigma)} = \frac{1}{2} \sum_{\sigma} L_1^{(\tau\sigma)} B_2^{(\tau\sigma)} \quad \left(\begin{matrix} \sigma = 1, 2, \dots, \mu_{\tau} \\ \varrho = 1, 2 \end{matrix} \right)$$

schreiben; dadurch erweist sich der Ausdruck $F_{12}^{(\tau)}$ als völlig unabhängig von der Wahl der Punkte $\mathfrak{f}_1^{(\tau\sigma)}$ und könnte daher nur noch mit der Lage von $\mathfrak{f}_2^{(\tau)}$ variieren.

Nun können wir bei beliebiger Wahl von $\mathfrak{f}_2^{(\tau)}$ in der Ebene $\{\mathfrak{F}_2^{(\tau)}\}$ die Punkte $\mathfrak{f}_i, \mathfrak{f}_i^{(\tau)}, \mathfrak{f}_i^{(\tau\sigma)}$ doch den sämtlichen Forderungen I. und II. gemäß einrichten. Alsdann sind zufolge I. die Werte $\xi^{(\tau\sigma\varrho)}, \eta^{(\tau\sigma)}$ für $\mathfrak{p}_1^{(\tau\sigma\varrho)} - \mathfrak{f}_1$ gleich $A_1^{(\tau\sigma\varrho)}, B_1^{(\tau\sigma)}$, für $\mathfrak{p}_2^{(\tau\sigma\varrho)} - \mathfrak{f}_2$ gleich $A_2^{(\tau\sigma\varrho)}, B_2^{(\tau\sigma)}$ und zufolge II. soll $\mathfrak{f}_1 = \mathfrak{f}_2 = \mathfrak{f}_0$ sein. Danach ergeben die senkrechten Projektionen der drei Punkte $\mathfrak{f}_0, \mathfrak{p}_1^{(\tau\sigma\varrho)}, \mathfrak{p}_2^{(\tau\sigma\varrho)}$ auf die Ebene $\zeta^{(\tau)} = 0$ in dieser ein Dreieck, bei welchem der Flächeninhalt gleich dem Betrage von

$$(114) \quad \frac{1}{2} (A_1^{(\tau\sigma\varrho)} B_2^{(\tau\sigma)} - A_2^{(\tau\sigma\varrho)} B_1^{(\tau\sigma)})$$

sein wird, während das Vorzeichen dieser Determinante nur dann das negative sein wird, wenn die hier angewiesene Folge der Ecken für dieses Dreieck einen entgegengesetzten Umlaufssinn bedeutet, als er sich bei einem Dreiecke in derselben Ebene $\zeta^{(\tau)} = 0$ einfindet, dessen drei Ecken nacheinander im Nullpunkte, auf der positiven $\xi^{(\tau\sigma\varrho)}$ -, auf der positiven $\eta^{(\tau\sigma)}$ -Achse angenommen werden.

Jetzt gehört zu jedem Indexsysteme τ, σ, ϱ nach den Ausführungen in § 19 ein bestimmtes zweites System $\tau, \sigma', \varrho' (\sigma' \neq \sigma)$, wobei stets $p_1^{(\tau\sigma\varrho)} = p_1^{(\tau\sigma'\varrho')}$, $p_2^{(\tau\sigma\varrho)} = p_2^{(\tau\sigma'\varrho')}$ ist, während für die beiden orthogonalen Substitutionen $S^{(\tau\sigma\varrho)}$ und $S^{(\tau\sigma'\varrho')}$ in ihrer gemeinsamen Koordinatenebene $\xi^{(\tau)} = 0$ das System der $\xi^{(\tau\sigma\varrho)}$ - und $\eta^{(\tau\sigma)}$ -Achse einerseits und das System der $\xi^{(\tau\sigma'\varrho')}$ - und $\eta^{(\tau\sigma')}$ -Achse entgegengesetzt orientiert erscheinen. Zufolge dieser Umstände muß der Ausdruck (114) nach seiner eben dargelegten Bedeutung bei Vertauschung von τ, σ, ϱ mit τ, σ', ϱ' einfach in den entgegengesetzten Wert übergehen, wir haben also stets

$$(115) \quad A_1^{(\tau\sigma\varrho)} B_2^{(\tau\sigma)} + A_1^{(\tau\sigma'\varrho')} B_2^{(\tau\sigma')} = A_2^{(\tau\sigma\varrho)} B_1^{(\tau\sigma)} + A_2^{(\tau\sigma'\varrho')} B_1^{(\tau\sigma')}.$$

Ordnen wir nun in dem Ausdrucke von $\mathfrak{F}_{12}^{(\tau)}$ die Glieder paarweise, indem wir mit jedem Gliede zu einem Indexsysteme τ, σ, ϱ das Glied zu dem korrespondierenden Indexsysteme τ, σ', ϱ' verbinden, so zeigt sich, daß

$$(116) \quad F_{12}^{(\tau)} = \frac{1}{2} \sum_{\sigma, \varrho} A_2^{(\tau\sigma\varrho)} B_1^{(\tau\sigma)} = F_{21}^{(\tau)} = \frac{1}{2} \sum_{\sigma} L_2^{(\tau\sigma)} B_1^{(\tau\sigma)}$$

ist.

Denken wir uns jetzt den Bezirk $\mathfrak{R}_2$ einer Translation parallel der Ebene $\{\mathfrak{F}_2^{(\tau)}\}$ unterworfen, wodurch $\mathfrak{F}_2^{(\tau)}$ dieselbe Translation in dieser Ebene erfährt, und halten wir dabei den Punkt $\mathfrak{f}_1 = \mathfrak{f}_2$ und die Beschränkungen I. und II. fest; alsdann bleiben alle Längen $L_2^{(\tau\sigma)}$ und die Werte $B_1^{(\tau\sigma)}$, $A_1^{(\tau\sigma\varrho)}$ ungeändert; es ändert sich also auch nicht $F_{21}^{(\tau)}$. Zugleich bleibt aber auch die Relation (116) und somit auch der Wert des Ausdrucks $F_{12}^{(\tau)}$ bestehen. Eine Translation von $\mathfrak{F}_2^{(\tau)}$ in der Ebene $\{\mathfrak{F}_2^{(\tau)}\}$ bei festem $\mathfrak{f}_2^{(\tau)}$ kommt nun für die Bildung des Ausdrucks (113) von $F_{12}^{(\tau)}$ auf das Gleiche hinaus wie die entgegengesetzte Translation des Punktes $\mathfrak{f}_2^{(\tau)}$ bei festgehaltenem Bereiche $\mathfrak{F}_2^{(\tau)}$. Danach ist endlich der Wert $F_{12}^{(\tau)}$ auch von der Lage von $\mathfrak{f}_2^{(\tau)}$ unabhängig. Zugleich ersehen wir, daß allgemein $F_{12}^{(\tau)} = F_{21}^{(\tau)}$ gilt und daß $F_{12}^{(\tau)}$ bei einer beliebigen Translation von $\mathfrak{F}_1^{(\tau)}$ oder von $\mathfrak{F}_2^{(\tau)}$ stets seinen Wert beibehält.

Die Größe $B_2^{(\tau\sigma)}$ ist die Differenz der $\eta^{(\tau\sigma)}$ -Werte für $p_2^{(\tau\sigma\varrho)}$ und für $\mathfrak{f}_2$, der zweite dieser Werte ist $= a_2 \alpha^{(\tau\sigma)} + b_2 \beta^{(\tau\sigma)} + c_2 \gamma^{(\tau\sigma)}$. Da nun der Wert $F_{12}^{(\tau)}$ nicht von der Lage von $\mathfrak{f}_2$ abhängt, so entnehmen wir aus (113) die drei Gleichungen:

$$(117) \quad \sum_{\sigma} \alpha^{(\tau\sigma)} L_1^{(\tau\sigma)} = 0, \quad \sum_{\sigma} \beta^{(\tau\sigma)} L_1^{(\tau\sigma)} = 0, \quad \sum_{\sigma} \gamma^{(\tau\sigma)} L_1^{(\tau\sigma)} = 0 \quad (\sigma = 1, 2, \dots, \mu_{\tau}).$$

Nehmen wir jetzt speziell für $\mathfrak{f}_2^{(\tau)}$, wenn $\mathfrak{F}_2^{(\tau)}$ ein Polygon ist, einen inneren Punkt dieses Polygons in seiner Ebene, wenn $\mathfrak{F}_2^{(\tau)}$ eine Strecke ist, einen von den Endpunkten verschiedenen Punkt dieser Strecke, wenn $\mathfrak{F}_2^{(\tau)}$ ein Punkt ist, eben diesen Punkt, so sind im ersten Falle alle Größen $B_2^{(\tau\sigma)} > 0$, im zweiten alle > 0 mit Ausnahme der beiden, die den zwei

Normalen der Strecke in der Ebene $\{\mathfrak{F}_2^{(\tau)}\}$ entsprechen und die $= 0$ sind, im dritten Falle alle $B_2^{(\tau\sigma)} = 0$. Wir entnehmen dann aus dem zweiten Ausdrucke in (113), daß $F_{12}^{(\tau)}$ stets ≥ 0 ausfällt und nur dann $= 0$ ist, wenn von den Bereichen $\mathfrak{F}_1^{(\tau)}$ und $\mathfrak{F}_2^{(\tau)}$ wenigstens einer ein Punkt ist oder diese beiden in zwei parallelen Geraden liegen. Wenn nicht dieser zweite Ausnahmefall statthat, ist überdies durch $\mathfrak{F}_1^{(\tau)}$ und $\mathfrak{F}_2^{(\tau)}$ die Ebene $\xi^{(\tau)} = 0$ völlig bestimmt, und da in (113) gewiß nur solchen Indizes σ ein von Null verschiedener Term entspricht, wobei $L_1^{(\tau\sigma)} > 0$ ist, die also wirklichen Kanten von $\mathfrak{F}_1^{(\tau)}$ entsprechen, so erkennen wir, daß der Wert von $F_{12}^{(\tau)}$ jedenfalls ganz allein von den zwei Bereichen $\mathfrak{F}_1^{(\tau)}$, $\mathfrak{F}_2^{(\tau)}$ und in keiner Weise von den sonst in Betracht kommenden Bereichen $\mathfrak{F}_3^{(\tau)}, \dots, \mathfrak{F}_m^{(\tau)}$ abhängt. Der Ausdruck $F_{11}^{(\tau)}$, in den $F_{12}^{(\tau)}$ übergeht, wenn $\mathfrak{F}_2^{(\tau)}$ mit $\mathfrak{F}_1^{(\tau)}$ identisch wird, bedeutet den Flächeninhalt von $\mathfrak{F}_1^{(\tau)}$. Wir wollen den Wert von $F_{12}^{(\tau)}$ als den *gemischten Flächeninhalt* von $\mathfrak{F}_1^{(\tau)}$ und $\mathfrak{F}_2^{(\tau)}$ bezeichnen.

Nun gilt

$$(118) \quad V_{123} = \frac{1}{6} \sum_{\tau, \sigma, \varrho} A_1^{(\tau\sigma\varrho)} B_2^{(\tau\sigma)} C_3^{(\tau)} = \frac{1}{3} \sum_{\tau} F_{12}^{(\tau)} C_3^{(\tau)},$$

und entsprechend können wir V_{213} darstellen; wegen (116) haben wir daher $V_{123} = V_{213}$.

Wir betrachten ferner die Transposition von 2 mit 3 im Ausdrucke V_{123} . Nach der Gleichung (118) und den letzten Resultaten in betreff der Größen $F_{12}^{(\tau)}$ könnte V_{123} jetzt nur noch von der Lage des Punktes $\mathfrak{f}_3$ abhängen. Wir können bei beliebiger Verfügung über $\mathfrak{f}_3$ doch die beschränkenden Annahmen I. und II. einführen. Alsdann sind die Koordinaten $\eta^{(\tau\sigma)}$, $\xi^{(\tau)}$ für $p_2^{(\tau\sigma\varrho)} - \mathfrak{f}_2$ gleich $B_2^{(\tau\sigma)}$, $C_2^{(\tau)}$, für $p_3^{(\tau\sigma\varrho)} - \mathfrak{f}_3$ gleich $B_3^{(\tau\sigma)}$, $C_3^{(\tau)}$. Zugleich soll $\mathfrak{f}_2 = \mathfrak{f}_3 = \mathfrak{f}_0$ sein, und besitzt daher das Dreieck, das durch senkrechte Projektion von $\mathfrak{f}_0 p_2^{(\tau\sigma\varrho)} p_3^{(\tau\sigma\varrho)}$ auf die Ebene $\xi^{(\tau\sigma\varrho)} = 0$ entsteht, einen Flächeninhalt gleich dem Betrage von

$$(119) \quad \frac{1}{2} (B_2^{(\tau\sigma)} C_3^{(\tau)} - B_3^{(\tau\sigma)} C_2^{(\tau)}),$$

während das Vorzeichen dieser Determinante dann das negative ist, wenn die hier angewiesene Folge der Ecken für dieses Dreieck einen entgegengesetzten Umlaufssinn bedeutet, als er sich für ein Dreieck in derselben Ebene $\xi^{(\tau\sigma\varrho)} = 0$ herausstellt, bei dem die Ecken nacheinander im Nullpunkte, auf der positiven $\eta^{(\tau\sigma)}$ -, auf der positiven $\xi^{(\tau)}$ -Achse angenommen sind.

Zu jedem Indexsysteme τ, σ, ϱ gehört nun nach § 19 ein bestimmtes zweites Indexsystem $\tau'', \sigma'', \varrho''$ ($\tau'' \neq \tau$), wobei immer $p_2^{(\tau\sigma\varrho)} = p_2^{(\tau''\sigma''\varrho'')}$, $p_3^{(\tau\sigma\varrho)} = p_3^{(\tau''\sigma''\varrho'')}$ und ferner $\xi^{(\tau\sigma\varrho)} = \xi^{(\tau''\sigma''\varrho'')}$ gilt. Die beiden Substitutionen $S^{(\tau\sigma\varrho)}$ und $S^{(\tau''\sigma''\varrho'')}$ haben dabei also die ξ -Achse gemein, dagegen sind in ihrer gemeinsamen Ebene $\xi^{(\tau\sigma\varrho)} = \xi^{(\tau''\sigma''\varrho'')} = 0$ das System der $\eta^{(\tau\sigma)}$ - und

$\xi^{(\tau)}$ -Achse einerseits und das System der $\eta^{(\tau''\sigma'')}$ - und $\xi^{(\tau'')}$ -Achse andererseits entgegengesetzt orientiert. Nach der eben angegebenen Bedeutung des Ausdrucks (119) entnehmen wir hieraus, daß dieser Ausdruck bei Ersetzung von τ, σ, ϱ durch $\tau'', \sigma'', \varrho''$ in den entgegengesetzten Wert übergeht, wir haben also stets

$$(120) \quad B_2^{(\tau\sigma)} C_3^{(\tau)} + B_2^{(\tau''\sigma'')} C_3^{(\tau'')} = B_3^{(\tau\sigma)} C_2^{(\tau)} + B_3^{(\tau''\sigma'')} C_2^{(\tau'')},$$

und zudem gilt

$$(121) \quad A_1^{(\tau\sigma\varrho)} = A_1^{(\tau''\sigma''\varrho'')}.$$

Ordnen wir jetzt die Glieder in V_{123} paarweise, indem wir mit jedem Gliede zu einem Indexsysteme τ, σ, ϱ das Glied zu dem zugehörigen Systeme $\tau'', \sigma'', \varrho''$ verbinden, so zeigt sich durch (120) und (121) unmittelbar, daß $V_{123} = V_{132}$ wird.

Unterwerfen wir nunmehr $\mathfrak{R}_3$ einer beliebigen Translation, während wir $\mathfrak{f}_3$ und die Annahmen I. und II. festhalten, so bleiben alle Größen im Ausdrucke

$$V_{132} = \frac{1}{3} \sum F_{13}^{(\tau)} C_2^{(\tau)}$$

und auch die letzte Relation bestehen und kann daher auch V_{123} sich nicht ändern. Eine Translation von $\mathfrak{R}_3$ bei festgehaltenem $\mathfrak{f}_3$ ist aber für die Bildung des Ausdrucks V_{123} in (118) gleichbedeutend mit der entgegengesetzten Translation von $\mathfrak{f}_3$ bei festem $\mathfrak{R}_3$, also hängt der Wert von V_{123} auch nicht von der Lage von $\mathfrak{f}_3$ ab. Zugleich ersehen wir, daß ganz allgemein $V_{123} = V_{132}$ gilt. Bezeichnet $H_3(u, v, w)$ die Stützebenenfunktion von $\mathfrak{R}_3$, so haben wir offenbar

$$(122) \quad H_3(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)}) = a_3 \alpha^{(\tau)} + b_3 \beta^{(\tau)} + c_3 \gamma^{(\tau)} + C_3^{(\tau)} \quad (\tau = 1, 2, \dots, \nu).$$

Da nun der Ausdruck (118) gar nicht von den Koordinaten a_3, b_3, c_3 abhängt, müssen danach stets die drei Beziehungen

$$(123) \quad \sum_{\tau} \alpha^{(\tau)} F_{12}^{(\tau)} = 0, \quad \sum_{\tau} \beta^{(\tau)} F_{12}^{(\tau)} = 0, \quad \sum_{\tau} \gamma^{(\tau)} F_{12}^{(\tau)} = 0 \quad (\tau = 1, 2, \dots, \nu)$$

bestehen, und insbesondere erhalten wir

$$(124) \quad V_{123} = \frac{1}{3} \sum_{\tau} F_{12}^{(\tau)} H_3(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)}).$$

Die beiden Relationen $V_{123} = V_{213}$ und $V_{123} = V_{132}$ zusammen bedeuten nun bereits die Tatsache, daß V_{123} bei beliebigen Permutationen der Indizes seinen Wert nicht verändert. Zugleich erkennen wir, daß V_{123} auch bei beliebigen Translationen der einzelnen drei Bezirke $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ stets ungeändert bleibt.

Wir denken uns jetzt $\mathfrak{f}_3$ speziell als einen *inwendigen* Punkt von $\mathfrak{R}_3$ angenommen, d. h. wenn $\mathfrak{R}_3$ ein Polyeder ist, soll $\mathfrak{f}_3$ ein innerer Punkt von $\mathfrak{R}_3$ sein, wenn $\mathfrak{R}_3$ ein Polygon oder eine Strecke ist, ein innerer

Punkt des Polygons in seiner Ebene bzw. der Strecke in ihrer Geraden, wenn $\mathfrak{R}_3$ ein Punkt ist, mit diesem Punkte identisch sein. Alsdann fallen die Größen $C_3^{(\tau)}$ sämtlich ≥ 0 aus und, soweit es dieser ihnen gemeinsame Charakter zuläßt, auch > 0 aus. Wir erkennen dann aus (118), daß V_{123} stets ≥ 0 ausfällt und insbesondere unter folgenden Umständen stets $= 0$ ist: 1. wenn unter den Bezirken $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ wenigstens ein Punkt oder 2. darunter wenigstens zwei Bezirke in parallelen Geraden vorhanden sind oder 3. diese Bezirke in drei parallelen Ebenen liegen. Einen von Null verschiedenen Term $F_{12}^{(\tau)} C_3^{(\tau)}$ haben wir in (118) nur, sowie für eine Richtung $(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)})$ einerseits $F_{12}^{(\tau)}$, andererseits $C_3^{(\tau)} > 0$ ist; alsdann sind $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ schon für sich jedenfalls Grundbezirke einer Polyederschar und kommt die Richtung bereits als Normale einer Seitenfläche dieser Schar vor. Danach hängt der Wert von V_{123} jedenfalls allein von $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ und in keiner Weise von den weiteren etwa betrachteten Bezirken $\mathfrak{R}_4, \dots, \mathfrak{R}_m$ ab. Die Größe V_{111} , in die V_{123} übergeht, wenn wir $\mathfrak{R}_2$ und $\mathfrak{R}_3$ mit $\mathfrak{R}_1$ identisch annehmen, bedeutet das Volumen von $\mathfrak{R}_1$. Wir wollen V_{123} das *gemischte Volumen* von $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ nennen und, wo es auf die deutlichere Bezeichnung der betreffenden Bezirke ankommt, auch $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ schreiben.

Nehmen wir für $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ beispielsweise drei nicht in einer Ebene gelegene Strecken op_1, op_2, op_3 mit o als Anfangspunkt, so können wir diese als drei Kanten für ein gewisses Parallelepipedum $\mathfrak{P}$ mit einer Ecke in o auffassen. Alsdann bedeutet $s_1 \mathfrak{R}_1 + s_2 \mathfrak{R}_2 + s_3 \mathfrak{R}_3$ für positive s_1, s_2, s_3 den Bereich desjenigen Parallelepipedum mit einer Ecke in o , bei dem die drei von o auslaufenden Kanten in den Punkten $s_1 p_1, s_2 p_2, s_3 p_3$ enden. Das Volumen dieses letzteren Parallelepipedum ist das $s_1 s_2 s_3$ -fache des Volumens von $\mathfrak{P}$, danach ist unter diesen Umständen $6 V_{123}$ gleich dem Volumen von $\mathfrak{P}$, mithin sicher $V_{123} > 0$.

Der Ausdruck $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ besitzt die folgende wichtige Eigenschaft:

Ist der konvexe Bezirk $\mathfrak{R}_3$ ganz enthalten in einem anderen konvexen Bezirk $\mathfrak{R}_3^*$ (ebenfalls mit einer endlichen Anzahl von extremen Punkten), so gilt stets

$$(125) \quad V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3) \leq V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3^*).$$

Denn bedeuten H_3 und H_3^* die Stützebenenfunktionen von $\mathfrak{R}_3$ und $\mathfrak{R}_3^*$, so gilt dann nach (44) allgemein $H_3(u, v, w) \leq H_3^*(u, v, w)$. Konstruieren wir nun eine Polyederschar, welche unter ihren Grundbezirken $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ und $\mathfrak{R}_3^*$ enthält, so bestehen gemäß (124) für $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ und $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3^*)$ gewisse Ausdrücke

$$\frac{1}{3} \sum_{\tau} F_{12}^{(\tau)} H_3(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)}), \quad \frac{1}{3} \sum_{\tau} F_{12}^{(\tau)} H_3^*(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)}),$$

und da hierin alle Faktoren $F_{12}^{(\tau)} \geq 0$ sind, geht in der Tat die Ungleichung (125) hervor.

Sind die drei konvexen Bezirke $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ derart beschaffen, daß darunter weder ein Punkt, noch zwei parallele Strecken vorhanden sind, noch diese drei Bezirke in drei parallelen Ebenen liegen, so können wir stets drei Strecken $\mathfrak{L}_1, \mathfrak{L}_2, \mathfrak{L}_3$ bestimmen, die nicht sämtlich einer Ebene parallel sind und so daß $\mathfrak{L}_1$ in $\mathfrak{R}_1$, $\mathfrak{L}_2$ in $\mathfrak{R}_2$, $\mathfrak{L}_3$ in $\mathfrak{R}_3$ liegt. Nach einer vorhin gemachten Bemerkung ist alsdann $V(\mathfrak{L}_1, \mathfrak{L}_2, \mathfrak{L}_3) > 0$, während dem letzten Satze zufolge sicher $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3) \geq V(\mathfrak{L}_1, \mathfrak{L}_2, \mathfrak{L}_3)$ wird. Also muß unter den angegebenen Umständen stets $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3) > 0$ ausfallen.

Aus (124) entnehmen wir ferner die Regeln:

Ist t ein positiver Parameter, so gilt

$$(126) \quad V(\mathfrak{R}_1, \mathfrak{R}_2, t\mathfrak{R}_3) = tV(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3).$$

Es besteht allgemein die Beziehung

$$(127) \quad V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3 + \mathfrak{R}_4) = V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3) + V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_4).$$

§ 22. Volumen in einer Schar konvexer Körper.

Es seien jetzt $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ m beliebige konvexe Bezirke, von denen ein jeder eine endliche oder unendliche Anzahl von extremen Punkten aufweisen kann. In jedem Bezirke $\mathfrak{R}_i$ ($i = 1, 2, \dots, m$) denken wir uns einen inwendigen Punkt e_i beliebig gewählt. Ferner bedeute ε irgendeine positive Größe. Nach dem Satze in § 5 und dem entsprechenden Satze für die Ebene können wir, wenn $\mathfrak{R}_i$ einen konvexen Körper oder ein ebenes Oval vorstellt, dazu ein Polyeder bzw. ein Polygon $\mathfrak{P}_i$ derart bestimmen, daß $\mathfrak{R}_i$ ganz in $\mathfrak{P}_i$ enthalten ist und andererseits selbst den Bereich $\frac{1}{1+\varepsilon}\mathfrak{P}_i + \frac{\varepsilon}{1+\varepsilon}e_i = \mathfrak{Q}_i$, der durch Dilatation des Bereichs $\mathfrak{P}_i$ von e_i aus im Verhältnis $\frac{1}{1+\varepsilon} : 1$ entsteht, ganz enthält. Bedeutet $\mathfrak{R}_i$ eine Strecke oder einen Punkt, so nehmen wir $\mathfrak{P}_i = \mathfrak{R}_i$ und haben noch eben diese Beziehungen zwischen $\mathfrak{R}_i, \mathfrak{P}_i, e_i, \varepsilon$. Die in solcher Art bestimmten Bereiche $\mathfrak{P}_i$ sind dann lauter konvexe Bezirke je mit einer endlichen Anzahl von extremen Punkten.

Bedeutet nun $s_1, s_2, \dots, s_m$ irgendwelche m Werte ≥ 0 , so wird jedesmal der Bezirk $\mathfrak{R} = \sum s_i \mathfrak{R}_i$ ganz im Bezirke $\mathfrak{P} = \sum s_i \mathfrak{P}_i$ enthalten sein, ferner den Punkt $e = \sum s_i e_i$ als inwendigen Punkt besitzen und den Bezirk $\frac{1}{1+\varepsilon}\mathfrak{P} + \frac{\varepsilon}{1+\varepsilon}e$ ganz in sich aufnehmen; dieser letztere Bezirk geht durch eine gewisse Translation aus $\frac{1}{1+\varepsilon}\mathfrak{P}$ hervor. Schreiben wir

$V(\mathfrak{P}_g, \mathfrak{P}_h, \mathfrak{P}_j) = W_{ghj}$, so finden wir dadurch für das Volumen V von $\mathfrak{R}$ die Ungleichungen

$$(128) \quad \sum W_{ghj} s_g s_h s_j \geq V \geq \frac{1}{(1+\varepsilon)^3} \sum W_{ghj} s_g s_h s_j \quad (g, h, j = 1, 2, \dots, m).$$

Wir bestimmen ferner eine positive Größe R so groß, daß der Würfel

$$|x| \leq R, \quad |y| \leq R, \quad |z| \leq R$$

alle Bereiche $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ völlig in sich aufnimmt. Alsdann enthält dieser Würfel auch jeden Bezirk $\mathfrak{Q}_i$, und mit Rücksicht auf den Satz bei (125) haben wir daher stets die Abschätzung:

$$(129) \quad \frac{1}{(1+\varepsilon)^3} W_{ghj} \leq 8R^3.$$

Jetzt denken wir uns noch auf irgendeine zweite Weise für jeden Bezirk $\mathfrak{R}_i$ einen inwendigen Punkt e_i^* ausgesucht und in bezug auf die nämliche Größe ε jedesmal einen konvexen Bezirk $\mathfrak{P}_i^*$ mit einer endlichen Anzahl von extremen Punkten irgendwie derart bestimmt, daß $\mathfrak{R}_i$ in $\mathfrak{P}_i^*$ enthalten ist und seinerseits den Bezirk $\frac{1}{1+\varepsilon} \mathfrak{P}_i^* + \frac{\varepsilon}{1+\varepsilon} e_i^* = \mathfrak{Q}_i^*$ enthält. Alsdann ist jedenfalls $\mathfrak{Q}_i^*$ in $\mathfrak{P}_i$ enthalten und haben wir, wenn $V(\mathfrak{P}_g^*, \mathfrak{P}_h^*, \mathfrak{P}_j^*) = W_{ghj}^*$ geschrieben wird, nach (125) stets

$$(130) \quad \frac{1}{(1+\varepsilon)^3} W_{ghj}^* \leq W_{ghj}.$$

Mit Benutzung von (129) erhalten wir daraus

$$(131) \quad W_{ghj}^* - W_{ghj} \leq ((1+\varepsilon)^3 - 1) (1+\varepsilon)^3 8R^3,$$

und hierin können wir noch W_{ghj} und W_{ghj}^* miteinander vertauschen, so daß die rechts stehende obere Grenze überhaupt für den Betrag der Differenz $W_{ghj}^* - W_{ghj}$ gilt. Daraus entnehmen wir, daß für ein nach Null abnehmendes ε jeder Ausdruck $V(\mathfrak{P}_g, \mathfrak{P}_h, \mathfrak{P}_j)$ nach einer bestimmten Grenze konvergiert, die wir mit $V(\mathfrak{R}_g, \mathfrak{R}_h, \mathfrak{R}_j)$ bezeichnen und *das gemischte Volumen* der konvexen Bezirke $\mathfrak{R}_g, \mathfrak{R}_h, \mathfrak{R}_j$ nennen. Aus den Ungleichungen (128) gewinnen wir alsdann, indem wir ε nach Null abnehmen lassen, für das Volumen V von $\mathfrak{R}$ allgemein die Entwicklung:

$$(132) \quad V = \sum V(\mathfrak{R}_g, \mathfrak{R}_h, \mathfrak{R}_j) s_g s_h s_j \quad (g, h, j = 1, 2, \dots, m).$$

Wir stellen jetzt die sämtlichen Eigenschaften des Ausdrucks $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ in bezug auf drei beliebige konvexe Bezirke $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ zusammen, die aus den entsprechenden Sätzen in § 21 unmittelbar folgen:

Der Wert von $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ ist unabhängig von der Reihenfolge der drei Bezirke; er ändert sich nicht bei beliebigen Translationen dieser einzelnen Bezirke.

Ist der Bezirk $\mathfrak{R}_3$ völlig enthalten in einem anderen konvexen Bezirk $\mathfrak{R}_3^*$, so gilt stets

$$(133) \quad V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3) \leq V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3^*).$$

In der Tat, wir können nach den oben gemachten Ausführungen stets vier konvexe Bezirke $\mathfrak{Q}_1, \mathfrak{Q}_2, \mathfrak{Q}_3, \mathfrak{P}_3^*$, einen jeden nur mit einer endlichen Anzahl von extremen Punkten, konstruieren, welche gewisse Annäherungen an $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_3^*$ vorstellen und zwar so, daß $\mathfrak{Q}_1$ in $\mathfrak{R}_1$, $\mathfrak{Q}_2$ in $\mathfrak{R}_2$, $\mathfrak{Q}_3$ in $\mathfrak{R}_3$ und andererseits $\mathfrak{R}_3^*$ in $\mathfrak{P}_3^*$ enthalten ist und daß dabei die Differenzen $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3) - V(\mathfrak{Q}_1, \mathfrak{Q}_2, \mathfrak{Q}_3)$ und $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3^*) - V(\mathfrak{Q}_1, \mathfrak{Q}_2, \mathfrak{P}_3^*)$ dem Betrage nach eine beliebig kleine positive GröÙe nicht übersteigen. Da nun auch $\mathfrak{Q}_3$ in $\mathfrak{P}_3^*$ enthalten ist und daher nach (125) sich

$$V(\mathfrak{Q}_1, \mathfrak{Q}_2, \mathfrak{P}_3^*) - V(\mathfrak{Q}_1, \mathfrak{Q}_2, \mathfrak{Q}_3) \geq 0$$

erweist, liegt dann $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3^*) - V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ über einer negativen GröÙe, deren Betrag wir beliebig klein einrichten können, d. h. diese letztere Differenz ist gleichfalls ≥ 0 .

Wir erkennen nun leicht, daß $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ stets ≥ 0 und nur in den Fällen $= 0$ ist, wenn 1. unter diesen drei Bezirken sich wenigstens ein Punkt oder 2. wenigstens zwei parallele Strecken finden oder 3. alle diese drei Bezirke ganz in drei parallelen Ebenen liegen.

Weiter übertragen sich die Regeln (126) und (127) sofort auf beliebige konvexe Bezirke. —

Für eine Schar mit nur zwei Grundbezirken $\mathfrak{R}, \mathfrak{R}'$ haben wir speziell die folgende Tatsache:

Sind $\mathfrak{R}$ und $\mathfrak{R}'$ irgend zwei konvexe Bezirke, so ist das Volumen eines Bezirks $s\mathfrak{R} + s'\mathfrak{R}'$, wobei $s \geq 0, s' \geq 0$ sind, durch einen kubischen Ausdruck

$$(134) \quad V_0 s^3 + 3V_1 s^2 s' + 3V_2 s s'^2 + V_3 s'^3$$

dargestellt. Darin bedeutet V_0 das Volumen von $\mathfrak{R}$, V_3 das Volumen von $\mathfrak{R}'$. Wir werden uns nun weiterhin eingehend mit den Eigenschaften des Koeffizienten $3V_1 = 3V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}')$ beschäftigen; bei Vertauschung der Rollen von $\mathfrak{R}$ und von $\mathfrak{R}'$ entsteht aus diesem Koeffizienten der Koeffizient $3V_2$ von ss'^2 .

In bezug auf $3V_1$ bemerken wir die folgende wichtige Darstellung:

Es bedeute $\mathfrak{P}$ ein konvexes Polyeder (oder Polygon) mit $\nu \geq 3$ (oder $\nu = 2$) Seitenflächen $\mathfrak{F}_1, \mathfrak{F}_2, \dots, \mathfrak{F}_\nu$. Es sei $(\alpha_\tau, \beta_\tau, \gamma_\tau)$ die äußere Normale und F_τ der Flächeninhalt von $\mathfrak{F}_\tau$ für $\tau = 1, 2, \dots, \nu$. Ferner sei $\mathfrak{R}'$ ein beliebiger konvexer Bezirk mit der Stützebenenfunktion $H'(u, v, w)$. Alsdann gilt die Darstellung:

$$(135) \quad 3V(\mathfrak{P}, \mathfrak{P}, \mathfrak{R}') = \sum_{\tau} F_{\tau} H'(\alpha_{\tau}, \beta_{\tau}, \gamma_{\tau}) \quad (\tau = 1, 2, \dots, \nu).$$

In der Tat, wir erhalten diese Beziehung sofort aus der Gleichung (124), wenn $\mathfrak{R}'$ eine endliche Anzahl von extremen Punkten besitzt. Insbesondere finden wir dadurch, indem wir für $\mathfrak{R}'$ einen beliebigen einzelnen Punkt a', b', c' nehmen, den Gleichungen (123) entsprechend, die drei Bedingungen:

$$(136) \quad \sum_{\tau} \alpha_{\tau} F_{\tau} = 0, \quad \sum_{\tau} \beta_{\tau} F_{\tau} = 0, \quad \sum_{\tau} \gamma_{\tau} F_{\tau} = 0 \quad (\tau = 1, 2, \dots, v).$$

Ist jedoch die Anzahl der extremen Punkte für $\mathfrak{R}'$ unendlich, so sei e' mit den Koordinaten a', b', c' irgendein inwendiger Punkt in $\mathfrak{R}'$. Als dann können wir, wenn eine positive Größe ε irgendwie angenommen wird, stets einen konvexen Bezirk $\mathfrak{P}'$ nur mit einer endlichen Anzahl von extremen Punkten derart bestimmen, daß $\mathfrak{R}'$ in $\mathfrak{P}'$ enthalten ist und selbst $\frac{1}{1+\varepsilon} \mathfrak{P}' + \frac{\varepsilon}{1+\varepsilon} e'$ enthält. Infolgedessen ist die Stützebenenfunktion von $\mathfrak{P}'$ für beliebige Argumente u, v, w

$$\geq H'(u, v, w) \text{ und } \leq (1 + \varepsilon) H'(u, v, w) - \varepsilon(a'u + b'v + c'w).$$

Nun können wir zur Berechnung von $3V(\mathfrak{P}, \mathfrak{P}, \mathfrak{P}')$ bereits die in (135) angewiesene Regel verwenden, und durch diese letzten Ungleichungen sowie unter Verwendung der Relationen (136) erhalten wir, da alle Faktoren $F_{\tau} > 0$ sind,

$$3V(\mathfrak{P}, \mathfrak{P}, \mathfrak{P}') \geq \sum_{\tau} F_{\tau} H'(\alpha_{\tau}, \beta_{\tau}, \gamma_{\tau}) \geq \frac{1}{1+\varepsilon} 3V(\mathfrak{P}, \mathfrak{P}, \mathfrak{P}').$$

In denselben zwei Grenzen wie die Summe hier, befindet sich aber, nach (133) auch die Größe $3V(\mathfrak{P}, \mathfrak{P}, \mathfrak{R}')$. Da nun diese Grenzen für ein nach Null abnehmendes ε einander beliebig nahe kommen, ist danach offenbar für $3V(\mathfrak{P}, \mathfrak{P}, \mathfrak{R}')$ genau die Formel (135) zutreffend.

Diese Gleichung (135) benutzen wir noch zu einer weiteren Folgerung:

Es bedeute $\mathfrak{R}$ einen konvexen Körper und es sei c_0 der kleinste, c_1 der größte Wert von z in $\mathfrak{R}$, ferner ξ ein beliebiger Wert $> c_0$ und $< c_1$. Die Ebene $z = \xi$ schneidet aus $\mathfrak{R}$ ein ebenes Oval heraus, das $\mathfrak{F}(\xi)$ und dessen Flächeninhalt $F(\xi)$ heiße; endlich bedeute $\mathfrak{R}_0$ denjenigen Teil des Körpers, in welchem $z \leq \xi$, und $\mathfrak{R}_1$ denjenigen Teil von $\mathfrak{R}$, in dem $z \geq \xi$ gilt; dabei sind $\mathfrak{R}_0$ und $\mathfrak{R}_1$ wieder zwei konvexe Körper. Ist nun $\mathfrak{R}'$ ein beliebiger konvexer Bezirk und $c'_0 (= -H'(0, 0, -1))$ der kleinste, $c'_1 (= H'(0, 0, 1))$ der größte Wert von z in $\mathfrak{R}'$, so haben wir stets die Gleichung:

$$(137) \quad V(\mathfrak{R}_0, \mathfrak{R}_0, \mathfrak{R}') + V(\mathfrak{R}_1, \mathfrak{R}_1, \mathfrak{R}') = V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}') + \frac{1}{3} (c'_1 - c'_0) F(\xi).$$

In der Tat, ist zunächst $\mathfrak{R}$ ein Polyeder, so sind auch $\mathfrak{R}_0$ und $\mathfrak{R}_1$ Polyeder und ist $\mathfrak{F}(\xi)$ ein Polygon. Dabei besitzt $\mathfrak{R}_0$ dieses Polygon als Seitenfläche mit der Normale $(0, 0, 1)$ und $\mathfrak{R}_1$ dieses Polygon als Seiten-

fläche mit der Normale $(0, 0, -1)$, während alle übrigen Seitenflächen von $\mathfrak{R}_0$ und $\mathfrak{R}_1$ zusammengenommen genau die vollständigen Seitenflächen von $\mathfrak{R}$ herstellen. Danach kommen wir nun sogleich zur Formel (137), wenn wir die drei darin genannten gemischten Volumina der Regel (135) gemäß ausdrücken. — Ist sodann $\mathfrak{R}$ ein konvexer Körper, so brauchen wir nur einen inneren Punkt e des Ovals $\mathfrak{F}(\xi)$ in seiner Ebene $z = \xi$ einzuführen, der dann zugleich ein innerer Punkt von $\mathfrak{R}$ wird, und, wenn ε eine beliebige positive Größe ist, ein Polyeder $\mathfrak{P}$ so zu bestimmen, daß $\mathfrak{R}$ in $\mathfrak{P}$ enthalten ist und selbst $\frac{1}{1+\varepsilon}\mathfrak{P} + \frac{\varepsilon}{1+\varepsilon}e$ enthält. Bilden wir dann die zu (137) entsprechende Gleichung ebenfalls in bezug auf die Schnittebene $z = \xi$ für ein solches Polyeder $\mathfrak{P}$, so folgt aus dieser durch Übergang zur Grenze $\lim \varepsilon = 0$ sofort die Relation (137) für $\mathfrak{R}$.

Zu den Sätzen dieses Abschnitts können wir vollkommen analoge Aussagen in der Geometrie der Ebene aufstellen. Zu je zwei konvexen Bezirken $\mathfrak{F}, \mathfrak{F}'$ in einer Ebene gehört stets ein bestimmter gemischter Flächeninhalt. Derselbe ist stets ≥ 0 und nur dann $= 0$, wenn wenigstens einer der Bezirke ein Punkt oder beide parallele Strecken sind; er hängt von der Reihenfolge der Bezirke nicht ab; er bleibt bei beliebigen Translationen der einzelnen Bezirke ungeändert; endlich wird er niemals verkleinert, wenn man einen der Bezirke durch einen umfassenderen konvexen Bezirk in der Ebene ersetzt. Sind $\mathfrak{F}_1, \mathfrak{F}_2, \dots, \mathfrak{F}_m$ m beliebige konvexe Bezirke in einer Ebene, so ist der Flächeninhalt eines Bezirks $\mathfrak{F} = \sum s_i \mathfrak{F}_i$ mit lauter Parametern $s_i \geq 0$ durch

$$(138) \quad F = \sum F_{gh} s_g s_h \quad (g, h = 1, 2, \dots, m)$$

dargestellt, wo F_{gh} den gemischten Flächeninhalt von $\mathfrak{F}_g$ und $\mathfrak{F}_h$ bedeutet.

§ 23. Der Schwerpunkt in einer Schar konvexer Körper.

Wir betrachten jetzt den Schwerpunkt eines Körpers in einer Schar konvexer Körper $\mathfrak{R} = \sum s_i \mathfrak{R}_i$ ($i = 1, 2, \dots, m; s_i \geq 0$).

Es handle sich zunächst um eine Polyederschar. Wir bezeichnen den Wert des Raumintegrals $\iiint x dx dy dz$ über den Bezirk $\mathfrak{R}$ mit $\mathfrak{X}$. Wie am Schlusse von § 20 ausgeführt ist, läßt sich $\mathfrak{X}$ als eine homogene Funktion vierten Grades der m Parameter s_i darstellen. Wir schreiben

$$(139) \quad \mathfrak{X} = \sum \mathfrak{X}_{ghjk} s_g s_h s_j s_k \quad (g, h, j, k = 1, 2, \dots, m),$$

wobei wir die Koeffizienten mit nicht lauter gleichen Indizes derart definieren, daß sie bei beliebigen Permutationen der vier Indizes sich nicht ändern sollen. Um zu zeigen, wie diese Entwicklung des Integrals $\mathfrak{X}$

sich auf Scharen mit völlig beliebigen Grundbezirken überträgt, haben wir mehrere Eigenschaften der Koeffizienten $\mathfrak{X}_{ghjk}$ abzuleiten.

Es wird genügen, die Bildungsweise eines Koeffizienten mit lauter verschiedenen Indizes ins Auge zu fassen, (wozu wir uns $m \geq 4$ eingerichtet denken). Die Größe $\mathfrak{X}_{1234}$ wird sich, nach den Gleichungen (101), (104) und (109), folgendermaßen entwickeln:

$$(140) \quad \mathfrak{X}_{1234} = \frac{1}{24 \cdot 24} \sum_{\tau, \sigma, \varrho} (a_g^{(\tau\sigma\varrho)} + a_g^{(\tau\sigma)} + a_g^{(\tau)} + a_g) A_h^{(\tau\sigma\varrho)} B_j^{(\tau\sigma)} C_k^{(\tau)},$$

worin für g, h, j, k nacheinander alle 24 verschiedenen Permutationen von 1, 2, 3, 4, einzusetzen sind und die Summation nach τ, σ, ϱ in der in § 20 erörterten Weise zu geschehen hat. Der Wert dieses Ausdrucks $\mathfrak{X}_{1234}$ ist nach seiner Bedeutung als Entwicklungskoeffizient in $\mathfrak{X}$ völlig unabhängig von den benutzten Hilfspunkten $\mathfrak{f}_i, \mathfrak{f}_i^{(\tau)}, \mathfrak{l}_i^{(\tau\sigma)}$ und hängt auch nur von den vier Bezirken $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4$ allein ab, falls die Zahl $m > 4$ ist; wir bezeichnen diese Größe auch in ausführlicherer Schreibweise mit $\mathfrak{X}(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4)$.

Die Stützebenenfunktion von $\mathfrak{R}_i$ heiße $H_i(u, v, w)$. Um eine gewisse Ungleichung abzuleiten, verfügen wir zunächst über die Hilfspunkte $\mathfrak{f}_i, \mathfrak{f}_i^{(\tau)}, \mathfrak{l}_i^{(\tau\sigma)}$ ($i = 1, 2, 3, 4$) in folgender besonderen Weise. Wir wählen jedesmal $\mathfrak{f}_i$ in $\mathfrak{R}_i$ selbst und zwar noch speziell in der Stützebene $x = H_i(1, 0, 0)$ an $\mathfrak{R}_i$, wo x den größten Wert in $\mathfrak{R}_i$ besitzt, ferner nehmen wir stets $\mathfrak{f}_i^{(\tau)}$ in $\mathfrak{R}_i^{(\tau)}$ selbst und $\mathfrak{l}_i^{(\tau\sigma)}$ in $\mathfrak{R}_i^{(\tau\sigma)}$ selbst an. Dann fallen alle Größen $A_i^{(\tau\sigma\varrho)}, B_i^{(\tau\sigma)}, C_i^{(\tau)} \geq 0$ aus, ferner wird das arithmetische Mittel von je drei Punkten $\mathfrak{p}_i^{(\tau\sigma\varrho)}, \mathfrak{l}_i^{(\tau\sigma)}, \mathfrak{f}_i^{(\tau)}$ stets ebenfalls noch ein Punkt in $\mathfrak{R}_i$ sein und für die x -Koordinate dieses Punktes dadurch sich jedesmal

$$-H_i(-1, 0, 0) \leq \frac{1}{3}(a_i^{(\tau\sigma\varrho)} + a_i^{(\tau\sigma)} + a_i^{(\tau)}) \leq H_i(1, 0, 0)$$

herausstellen, und endlich haben wir noch $a_i = H_i(1, 0, 0)$. Mit Rücksicht auf diese Umstände gewinnen wir aus (140) die Ungleichungen:

$$(141) \quad \begin{aligned} \frac{1}{4} \sum_{g=1}^4 H_g(1, 0, 0) V_{ghjk} &\geq \mathfrak{X}_{1234} \geq \\ &\geq \frac{1}{16} \sum_{g=1}^4 (-3 H_g(-1, 0, 0) + H_g(1, 0, 0)) V_{ghjk}, \end{aligned}$$

wobei für g nacheinander 1, 2, 3, 4 zu nehmen ist, für h, j, k jedesmal die drei von g verschiedenen der Indizes 1, 2, 3, 4 nach der Größe geordnet zu setzen sind und V_{ghjk} das gemischte Volumen von $\mathfrak{R}_h, \mathfrak{R}_j, \mathfrak{R}_k$ gemäß der Formel (108) bedeutet.

Nehmen wir die Bezirke $\mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4$ als identisch mit $\mathfrak{R}_1$ an, so geht daraus spezieller

$$(142) \quad H_1(1, 0, 0) V_{111} \geq \mathfrak{X}_{111} \geq \frac{1}{4} (-3 H_1(-1, 0, 0) + H_1(1, 0, 0)) V_{111}$$

hervor; darin bedeutet dann V_{111} das Volumen von $\mathfrak{R}_1$. Ist nun $V_{111} > 0$, also $\mathfrak{R}_1$ ein konvexer Körper und denken wir uns den Nullpunkt der Koordinaten mit dem Schwerpunkte dieses Körpers zusammenfallend, so haben wir $\mathfrak{X}_{111} = 0$ und finden $3 H_1(-1, 0, 0) \geq H_1(1, 0, 0)$. Da wir für die Richtung der x -Achse eine ganz beliebige Richtung einführen können, so gilt, wenn der Schwerpunkt von $\mathfrak{R}_1$ im Nullpunkte liegt, überhaupt allgemein für jedes Paar entgegengesetzter Richtungen $(-\alpha, -\beta, -\gamma)$ und (α, β, γ) :

$$(143) \quad 2 H_1(\alpha, \beta, \gamma) \geq \frac{1}{2} (H_1(\alpha, \beta, \gamma) + H_1(-\alpha, -\beta, -\gamma)).$$

Die hier rechts stehende Größe bedeutet den halben gegenseitigen Abstand der zwei Stützebenen an $\mathfrak{R}_1$ mit den Normalen (α, β, γ) und $(-\alpha, -\beta, -\gamma)$ und ist niemals kleiner als der Radius der größten überhaupt in $\mathfrak{R}_1$ enthaltenen Kugel, während das Minimum von $H_1(\alpha, \beta, \gamma)$ gleich dem Radius der größten in $\mathfrak{R}_1$ enthaltenen Kugel mit dem Schwerpunkte von $\mathfrak{R}_1$ als Mittelpunkt ist. Danach ist der letztere Radius mindestens halb so groß wie der an erster Stelle genannte Radius, eine bereits in § 9 ohne Beweis angeführte Tatsache.

Wir wollen jetzt die Abhängigkeit des Ausdrucks $\mathfrak{X}(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4)$ von einem einzelnen der betreffenden vier Bezirke, etwa von $\mathfrak{R}_4$, ins Auge fassen. Wir werden zeigen, daß wir aus dem Ausdrucke (140) die auf $\mathfrak{R}_4$ bezüglichen Größen $A_4^{(\tau\sigma\varrho)}$ und $B_4^{(\tau\sigma)}$ eliminieren können, so daß wir darin nur noch die Größen $C_4^{(\tau)}$ behalten, für welche wir

$$(144) \quad C_4^{(\tau)} = -a_4 \alpha^{(\tau)} - b_4 \beta^{(\tau)} - c_4 \gamma^{(\tau)} + H_4(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)})$$

haben, wenn a_4, b_4, c_4 die Koordinaten des Punktes $\mathfrak{f}_4$ sind.

Zum Zwecke dieser Umformung von (140) unterwerfen wir die Hilfspunkte $\mathfrak{f}_i, \mathfrak{f}_i^{(\tau)}, \mathfrak{l}_i^{(\tau\sigma)}$ ($i = 1, 2, 3, 4$) zunächst wieder den beschränkenden Annahmen I. und II. aus § 21. Es sei also jedesmal $\mathfrak{f}_i^{(\tau)}$ der Fußpunkt des von $\mathfrak{f}_i$ auf die Ebene $\{\mathfrak{F}_i^{(\tau)}\}$, $\mathfrak{l}_i^{(\tau\sigma)}$ der Fußpunkt des von $\mathfrak{f}_i$ auf die Gerade $\{\mathfrak{L}_i^{(\tau\sigma)}\}$ gefällten Lotes, und ferner sei $\mathfrak{f}_1 = \mathfrak{f}_2 = \mathfrak{f}_3 = \mathfrak{f}_4$ angenommen, während die Koordinaten des einen Punktes $\mathfrak{f}_4$ noch völlig beliebig bleiben. Nach (110) haben wir

$$a_g^{(\tau\sigma\varrho)} + a_g^{(\tau\sigma)} + a_g^{(\tau)} + a_g = \alpha^{(\tau\sigma\varrho)} A_g^{(\tau\sigma\varrho)} + 2 \alpha^{(\tau\sigma)} B_g^{(\tau\sigma)} + 3 \alpha^{(\tau)} C_g^{(\tau)} + 4 a_g.$$

Indem wir diesen Ausdruck in (140) einführen und den Index 4 besonders hervorheben, können wir die Differenz $\mathfrak{X}_{1234} - \frac{1}{4} V_{123} a_4$ zunächst in folgender Weise darstellen:

$$\begin{aligned}
 (145) \quad & \frac{1}{24^2} \sum_{\tau, \sigma, \varrho} \sum (2\alpha^{(\tau\sigma\varrho)} A_g^{(\tau\sigma\varrho)} + 2\alpha^{(\tau\sigma)} B_g^{(\tau\sigma)} + 3\alpha^{(\tau)} C_g^{(\tau)} + 4a_g) A_4^{(\tau\sigma\varrho)} B_h^{(\tau\sigma)} C_j^{(\tau)} \\
 & + \frac{1}{24^2} \sum_{\tau, \sigma, \varrho} \sum (\alpha^{(\tau\sigma\varrho)} A_g^{(\tau\sigma\varrho)} + 4\alpha^{(\tau\sigma)} B_g^{(\tau\sigma)} + 3\alpha^{(\tau)} C_g^{(\tau)} + 4a_g) A_h^{(\tau\sigma\varrho)} B_4^{(\tau\sigma)} C_j^{(\tau)} \\
 & + \frac{1}{24^2} \sum_{\tau, \sigma, \varrho} \sum (\alpha^{(\tau\sigma\varrho)} A_g^{(\tau\sigma\varrho)} + 2\alpha^{(\tau\sigma)} B_g^{(\tau\sigma)} + 6\alpha^{(\tau)} C_g^{(\tau)} + 4a_g) A_h^{(\tau\sigma\varrho)} B_j^{(\tau\sigma)} C_4^{(\tau)},
 \end{aligned}$$

wobei dann g, h, j alle 6 Permutationen von 1, 2, 3 zu durchlaufen haben und die Summation über τ, σ, ϱ in der bisherigen Weise zu nehmen ist. Der in der ersten Reihe eingeklammerte Ausdruck ist $= 2a_g^{(\tau\sigma\varrho)} + a_g^{(\tau)} + a_g$. Nach den Ausführungen in § 21 gehört mit jedem Indexsysteme τ, σ, ϱ ein bestimmtes zweites System τ', σ', ϱ' ($\sigma' \neq \sigma$) zusammen, derart, daß dabei stets (s. (115))

$$A_4^{(\tau\sigma\varrho)} B_h^{(\tau\sigma)} + A_4^{(\tau\sigma'\varrho')} B_h^{(\tau\sigma')} = A_h^{(\tau\sigma\varrho)} B_4^{(\tau\sigma)} + A_h^{(\tau\sigma'\varrho')} B_4^{(\tau\sigma')}$$

und zugleich $a_g^{(\tau\sigma\varrho)} = a_g^{(\tau\sigma'\varrho')}$ gilt. Indem wir die Glieder bei der ersten Summation $\sum_{\tau, \sigma, \varrho}$ in (145) dementsprechend paarweise zusammenfassen, erkennen wir hieraus, daß wir in jenem ersten Aggregat, ohne seinen Wert zu ändern, die Indizes 4 und h miteinander vertauschen dürfen. Der Ausdruck (145) verwandelt sich dadurch in

$$\begin{aligned}
 (146) \quad & \frac{1}{24^2} \sum_{\tau, \sigma, \varrho} \sum (3\alpha^{(\tau\sigma\varrho)} A_g^{(\tau\sigma\varrho)} + 6\alpha^{(\tau\sigma)} B_g^{(\tau\sigma)} + 6\alpha^{(\tau)} C_g^{(\tau)} + 8a_g) A_h^{(\tau\sigma\varrho)} B_4^{(\tau\sigma)} C_j^{(\tau)} \\
 & + \frac{1}{24^2} \sum_{\tau, \sigma, \varrho} \sum (\alpha^{(\tau\sigma\varrho)} A_g^{(\tau\sigma\varrho)} + 2\alpha^{(\tau\sigma)} B_g^{(\tau\sigma)} + 6\alpha^{(\tau)} C_g^{(\tau)} + 4a_g) A_h^{(\tau\sigma\varrho)} B_j^{(\tau\sigma)} C_4^{(\tau)}.
 \end{aligned}$$

Hier ist der Faktor in der ersten Klammer $= 3a_g^{(\tau\sigma\varrho)} + 3a_g^{(\tau\sigma)} + 2a_g$. Weiter gehört nach § 21 mit jedem System τ, σ, ϱ ein bestimmtes zweites System $\tau'', \sigma'', \varrho''$, ($\tau'' \neq \tau$) zusammen, derart, daß dabei stets (s. (120))

$$B_4^{(\tau\sigma)} C_j^{(\tau)} + B_4^{(\tau''\sigma'')} C_j^{(\tau'')} = B_j^{(\tau\sigma)} C_4^{(\tau)} + B_j^{(\tau''\sigma'')} C_4^{(\tau'')}$$

und zugleich $a_g^{(\tau\sigma\varrho)} = a_g^{(\tau''\sigma''\varrho'')}$, $a_g^{(\tau\sigma)} = a_g^{(\tau''\sigma'')}$, $A_h^{(\tau\sigma\varrho)} = A_h^{(\tau''\sigma''\varrho'')}$ gilt. Danach können wir weiter in dem ersten Aggregat in (146), ohne seinen Wert zu ändern, die Indizes 4 und j miteinander vertauschen; der ganze Ausdruck (146) geht dadurch in

$$(147) \quad \frac{1}{24^2} \sum_{\tau, \sigma, \varrho} \sum (4\alpha^{(\tau\sigma\varrho)} A_g^{(\tau\sigma\varrho)} + 8\alpha^{(\tau\sigma)} B_g^{(\tau\sigma)} + 12\alpha^{(\tau)} C_g^{(\tau)} + 12a_g) A_h^{(\tau\sigma\varrho)} B_j^{(\tau\sigma)} C_4^{(\tau)}$$

über, und erlangen wir damit endlich die Formel

$$(148) \quad \mathfrak{X}_{1234} = \frac{1}{4} V_{123} \alpha_4 + \frac{1}{24 \cdot 6} \sum_{\tau, \sigma, \varrho} \sum (a_g^{(\tau\sigma\varrho)} + a_g^{(\tau\sigma)} + a_g^{(\tau)}) A_h^{(\tau\sigma\varrho)} B_j^{(\tau\sigma)} C_4^{(\tau)},$$

wo wieder g, h, j alle 6 Permutationen von 1, 2, 3 zu durchlaufen haben

Wir können nun den hier rechts stehenden Ausdruck auch bilden, ohne über die Hilfspunkte die beschränkenden Voraussetzungen I. und II. zu machen, und wir finden alsdann diesen Ausdruck stets von demselben Wert. In der Tat, denken wir uns jetzt die Hilfspunkte $\mathfrak{f}_i$, $\mathfrak{f}_i^{(\tau)}$, $\mathfrak{l}_i^{(\tau\sigma)}$ wieder in der vollen Allgemeinheit wie früher angenommen. Führen wir anstatt x, y, z neue orthogonale Koordinaten ξ, η, ζ auf irgendeine Weise derart ein, daß die positive ξ -Achse die Richtung $(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)})$ hat, so können wir das Flächenintegral $\mathfrak{X}^{(\tau)} = \iint x d\xi d\eta$ über das Gebiet $\mathfrak{F}^{(\tau)}$ in der Stützebene $\{\mathfrak{F}^{(\tau)}\}$ des Bezirks $\mathfrak{R}$, (ähnlich wie wir für $\mathfrak{X}$ den Ausdruck (101) herstellten), in der Formel

$$(149) \quad \mathfrak{X}^{(\tau)} = \frac{1}{6} \sum_{\sigma, \varrho} (a^{(\tau\sigma\varrho)} + a^{(\tau\sigma)} + a^{(\tau)}) A^{(\tau\sigma\varrho)} B^{(\tau\sigma)} \quad \left(\begin{matrix} \sigma = 1, 2, \dots, \mu_\tau; \\ \varrho = 1, 2 \end{matrix} \right)$$

entwickeln. Vermöge der Gleichungen in (104) und (109) erweist sich dann dieser Wert $\mathfrak{X}^{(\tau)}$ als eine kubische Form in $s_1, s_2, \dots, s_m$. Wir schreiben dieselbe

$$(150) \quad \mathfrak{X}^{(\tau)} = \sum \mathfrak{X}_{ghj}^{(\tau)} s_g s_h s_j \quad (g, h, j = 1, 2, \dots, m),$$

so daß $\mathfrak{X}_{ghj}^{(\tau)}$ bei einer Permutation der Indizes g, h, j sich nicht verändert. Alsdann haben wir insbesondere

$$(151) \quad \mathfrak{X}_{123}^{(\tau)} = \frac{1}{6^2} \sum_{\sigma, \varrho} (a_g^{(\tau\sigma\varrho)} + a_g^{(\tau\sigma)} + a_g^{(\tau)}) A_h^{(\tau\sigma\varrho)} B_j^{(\tau\sigma)} \quad \left(\begin{matrix} \sigma = 1, 2, \dots, \mu_\tau; \\ \varrho = 1, 2 \end{matrix} \right),$$

und kommt danach die Formel (148) auf die Entwicklung

$$(152) \quad \mathfrak{X}_{1234} = \frac{1}{4} V_{123} a_4 + \frac{1}{4} \sum_{\tau} \mathfrak{X}_{123}^{(\tau)} C_4^{(\tau)} \quad (\tau = 1, 2, \dots, \nu)$$

hinaus.

Jede Größe $\mathfrak{X}_{123}^{(\tau)}$ hier ist nach ihrer Bedeutung als Entwicklungskoeffizient im Ausdrucke (150) von $\mathfrak{X}^{(\tau)}$ durchaus unabhängig von den benutzten Hilfspunkten. Nehmen wir jetzt insbesondere für $i = 1, 2, 3$ immer $\mathfrak{f}_i^{(\tau)}$ in $\mathfrak{F}_i^{(\tau)}$ selbst, $\mathfrak{l}_i^{(\tau\sigma)}$ in $\mathfrak{L}_i^{(\tau\sigma)}$ selbst, so sind in (151) die Größen $A_i^{(\tau\sigma\varrho)}$, $B_i^{(\tau\sigma)}$ sämtlich ≥ 0 . Wir denken uns ferner eine positive Größe R derart bestimmt, daß $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4$ ganz im Würfel

$$|x| \leq R, \quad |y| \leq R, \quad |z| \leq R$$

enthalten sind; alsdann ist weiter jede Größe $a_i^{(\tau\sigma\varrho)}$, $a_i^{(\tau\sigma)}$, $a_i^{(\tau)}$ in (151) dem Betrage nach $\leq R$ und finden wir daher

$$(153) \quad |\mathfrak{X}_{123}^{(\tau)}| \leq \frac{1}{3} R (F_{12}^{(\tau)} + F_{13}^{(\tau)} + F_{23}^{(\tau)}), \quad V_{123} \leq 8 R^3.$$

Die Relation (152) benutzen wir, um die Veränderung abzuschätzen, welche der Wert $\mathfrak{X}(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4)$ bei einer Veränderung der Grundbezirke $\mathfrak{R}_i$ erfährt. Wir nehmen in jedem Bezirke $\mathfrak{R}_i$ ($i = 1, 2, 3, 4$) einen in-

wendigen Punkt e_i beliebig an, ferner sei ε eine beliebige positive Größe und sodann $\mathfrak{R}_i^*$ jedesmal ein solcher konvexer Bezirk ebenfalls mit einer endlichen Anzahl von extremen Punkten, der ganz in $\mathfrak{R}_i$ enthalten ist und seinerseits ganz den Bereich $\frac{1}{1+\varepsilon} \mathfrak{R}_i + \frac{\varepsilon}{1+\varepsilon} e_i$ enthält. Wir denken uns nun eine Polyederschar gebildet, welche sämtliche Bezirke $\mathfrak{R}_i, \mathfrak{R}_i^* (i=1, 2, 3, 4)$ unter ihren Grundbezirken enthält, und können dann zur Berechnung von $\mathfrak{X}(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4)$ und andererseits von $\mathfrak{X}(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4^*)$ die Regel (152) mit Bezug auf diese Schar verwenden, so daß also darin die Summation über alle solchen Richtungen $(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)})$ auszudehnen ist, welche als Normalen der Seitenflächen bei dieser weiteren Polyederschar auftreten. Legen wir nun den Punkt f_4 sowie den entsprechenden Hilfspunkt f_4^* für $\mathfrak{R}_4^*$ in den Punkt e_4 , so ist jede Größe $C_4^{(\tau)} \geq 0$ und die zu $C_4^{(\tau)}$ entsprechende Größe $C_4^{*(\tau)}$ für $\mathfrak{R}_4^*$ jedesmal $\leq C_4^{(\tau)}$ und $\geq \frac{1}{1+\varepsilon} C_4^{(\tau)}$, und finden wir sodann mit Rücksicht auf (152) und (153):

$$(154) \quad |\mathfrak{X}(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4) - \mathfrak{X}(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4^*)| \leq \frac{1}{4} \frac{\varepsilon}{1+\varepsilon} R(V_{124} + V_{134} + V_{234}) \leq \frac{6\varepsilon}{1+\varepsilon} R^4.$$

Von dieser Ungleichung aus kommen wir sogleich weiter zu der Beziehung:

$$(155) \quad |\mathfrak{X}(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4) - \mathfrak{X}(\mathfrak{R}_1^*, \mathfrak{R}_2^*, \mathfrak{R}_3^*, \mathfrak{R}_4^*)| \leq \frac{24\varepsilon}{1+\varepsilon} R^4.$$

Auf diese Abschätzung fußend können wir, ähnlich wie wir in § 22 den Begriff des gemischten Volumens erweiterten, jetzt auch den Begriff der Bildung $\mathfrak{X}(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3, \mathfrak{R}_4)$ auf beliebige konvexe Bezirke ausdehnen, die nicht bloß eine endliche Anzahl von extremen Punkten besitzen, und dabei übertragen sich dann die Ungleichungen (141) und (155) sofort auch auf beliebige konvexe Bezirke. Wir erhalten das Resultat, daß für jede beliebige Schar konvexer Bezirke $\mathfrak{R} = \sum s_i \mathfrak{R}_i (i=1, 2, \dots, m; s_i \geq 0)$ das Raumintegral $\mathfrak{X} = \iiint x dx dy dz$ über $\mathfrak{R}$ sich immer in der Form

$$(156) \quad \mathfrak{X} = \sum \mathfrak{X}(\mathfrak{R}_g, \mathfrak{R}_h, \mathfrak{R}_j, \mathfrak{R}_k) s_g s_h s_j s_k \quad (g, h, j, k = 1, 2, \dots, m)$$

entwickeln läßt.

IV. Kapitel.

Zweigliedrige Scharen.

§ 24. Definition der Oberfläche eines konvexen Körpers.

Bei einem konvexen *Polyeder* erklären wir als die *Oberfläche* des Polyeders die Summe der Flächeninhalte seiner einzelnen Seitenflächen.

Wir bezeichnen wie schon früher mit $\mathfrak{G}$ die Kugel vom Radius 1 mit dem Nullpunkte o als Mittelpunkt. Die Stützebenenfunktion dieser

Kugel $\mathfrak{G}$ ist für solche Argumente α, β, γ , welche der Bedingung $\alpha^2 + \beta^2 + \gamma^2 = 1$ genügen, konstant $= 1$. Vermöge der Formel (135) erkennen wir daher:

Die Oberfläche eines konvexen Polyeders $\mathfrak{P}$ ist stets gleich dem Dreifachen des gemischten Volumens von $\mathfrak{P}, \mathfrak{P}, \mathfrak{G}$.

Jetzt definieren wir allgemein die Oberfläche eines beliebigen konvexen Körpers $\mathfrak{R}$ als das Dreifache des gemischten Volumens von $\mathfrak{R}, \mathfrak{R}, \mathfrak{G}$, also ihre Größe $\mathfrak{D}$ durch die Gleichung

$$(157) \quad \mathfrak{D} = 3 V(\mathfrak{R}, \mathfrak{R}, \mathfrak{G}).$$

Aus dem Satze in (133) entnehmen wir dann insbesondere die Folgerung:

Ist ein konvexer Körper $\mathfrak{R}$ ganz in einem anderen konvexen Körper $\mathfrak{R}^*$ enthalten, so ist die Oberfläche $\mathfrak{D}$ von $\mathfrak{R}$ gewiß niemals kleiner als die Oberfläche $\mathfrak{D}^*$ von $\mathfrak{R}^*$. Wir werden sogleich (in § 26 und § 27) sehen, daß dabei sogar immer $\mathfrak{D} < \mathfrak{D}^*$ gilt, also niemals $\mathfrak{D}$ gleich $\mathfrak{D}^*$ sein kann.

Analog wie wir hier für jeden konvexen Körper eine bestimmte „Oberfläche“ einführen, definieren wir für jedes ebene Oval $\mathfrak{F}$ einen bestimmten „Umfang“ von $\mathfrak{F}$ als das Doppelte des gemischten Flächeninhalts von $\mathfrak{F}$ und einem Kreise mit dem Radius 1 in der Ebene von $\mathfrak{F}$.

Bedeutet $\mathfrak{R}$ einen bestimmten konvexen Körper, t einen beliebigen positiven Parameter, so ist das Volumen des konvexen Körpers $\mathfrak{R} + t\mathfrak{G}$ gemäß (134) durch einen Ausdruck

$$(158) \quad V_0 + 3 V_1 t + 3 V_2 t^2 + V_3 t^3$$

darzustellen; darin bezeichnet nun V_0 das Volumen, $3V_1$ die Oberfläche von $\mathfrak{R}$, weiter V_2 das gemischte Volumen $V(\mathfrak{R}, \mathfrak{G}, \mathfrak{G})$ und endlich V_3 das Volumen von $\mathfrak{G}$, also die Konstante $\frac{4\pi}{3}$. Fassen wir die Begrenzung von $\mathfrak{R} + t\mathfrak{G}$ näher ins Auge. Nach dem allgemeinen Satze in § 18 läßt jeder Punkt q der Begrenzung des Körpers $\mathfrak{R} + t\mathfrak{G}$ eine Darstellung $q = p + tg$ zu, so daß p ein Punkt von $\mathfrak{R}$ und g ein Punkt von $\mathfrak{G}$ ist, und zwar muß dann notwendig p der Begrenzung von $\mathfrak{R}$ und g der Begrenzung von $\mathfrak{G}$ angehören und müssen zudem zwei Stützebenen mit der nämlichen äußeren Normalen durch p an $\mathfrak{R}$ und durch g an $\mathfrak{G}$ gehen. Da nun durch jeden Punkt g der Begrenzung der Kugel $\mathfrak{G}$ stets nur eine einzige Stützebene an $\mathfrak{G}$, nämlich die Tangentialebene durch g an $\mathfrak{G}$ existiert, so stellt danach die Begrenzung von $\mathfrak{R} + t\mathfrak{G}$ offenbar nichts anderes vor als die zur Begrenzung von $\mathfrak{R}$ im konstanten Normalabstande t nach dem Äußeren von $\mathfrak{R}$ hin konstruierte *Parallelfläche*.

Auf diese Bemerkung gestützt können wir das Zustandekommen der Formel (158) auf einem direkten Wege einsehen im Falle, daß $\mathfrak{R}$ speziell

eine endliche Anzahl von extremen Punkten aufweist. Alsdann können wir nämlich den ganzen Körper $\mathfrak{R} + t\mathfrak{G}$ aus dem Polyeder $\mathfrak{R}$ als Kern unter Hinzufügung einer endlichen Anzahl einfach konstruierter und untereinander völlig getrennt liegender Zusatzstücke aufbauen.

Wir wollen die Seitenflächen, Kanten, Ecken des Polyeders $\mathfrak{R}$ wie in § 19 bezeichnen. Um das Polyeder $\mathfrak{R}$ zum Körper $\mathfrak{R} + t\mathfrak{G}$ auszugestalten, haben wir *erstens* auf jede Seitenfläche $\mathfrak{F}^{(\tau)}$ von $\mathfrak{R}$ als Grundfläche nach dem Äußeren von $\mathfrak{R}$ hin ein senkrechtes *Prisma* mit einer Höhe von der Länge t aufzubauen, also vom Volumen $tF^{(\tau)}$, unter $F^{(\tau)}$ den Flächeninhalt von $\mathfrak{F}^{(\tau)}$ verstanden. Das Prisma über $\mathfrak{F}^{(\tau)}$ besitzt über jeder Kante $\mathfrak{Q}^{(\tau\sigma)}$ von $\mathfrak{F}^{(\tau)}$ eine Seitenfläche in Form eines Rechtecks mit der Basis $\mathfrak{Q}^{(\tau\sigma)}$ und einer Höhe $= t$.

An einer beliebigen Kante $\mathfrak{Q}^{(\tau\sigma)}$ von $\mathfrak{R}$ stoßen nunmehr immer zwei derartige, einander kongruente rechteckige Seitenflächen von zwei jener Prismen zusammen, des über $\mathfrak{F}^{(\tau)}$ und des über $\mathfrak{F}^{(\tau')}$, wenn $\mathfrak{Q}^{(\tau\sigma)} = \mathfrak{Q}^{(\tau''\sigma')}$ wie in § 19 gedacht ist. Alsdann haben wir *zweitens* jedesmal an der Kante $\mathfrak{Q}^{(\tau\sigma)}$ als Achse denjenigen *Zylinderausschnitt* zu konstruieren, der durch eine im Äußeren von $\mathfrak{R}$ erfolgende Rotation des einen Rechtecks um die gemeinsame Grundlinie $\mathfrak{Q}^{(\tau\sigma)}$ bis zum Zusammenfallen mit dem anderen Rechteck entsteht. Ist $L^{(\tau\sigma)}$ die Länge von $\mathfrak{Q}^{(\tau\sigma)}$ und $\chi^{(\tau\sigma)}$ der Winkel, den die äußeren Normalen der beiden Seitenflächen $\mathfrak{F}^{(\tau)}$ und $\mathfrak{F}^{(\tau')}$ miteinander bilden, so ist $\frac{1}{2}t^2\chi^{(\tau\sigma)}L^{(\tau\sigma)}$ das Volumen des auf solche Weise beschriebenen Körpers. Dieser Zylinderausschnitt an $\mathfrak{Q}^{(\tau\sigma)}$ besitzt an jeder Ecke $\mathfrak{p}^{(\tau\sigma\varrho)}$ von $\mathfrak{Q}^{(\tau\sigma)}$ eine Grundfläche in Form eines Sektors aus einer Kreisfläche vom Radius t .

An einer beliebigen Ecke $\mathfrak{p}^{(\tau\sigma\varrho)}$ von $\mathfrak{R}$ stoßen nunmehr immer genau so viele derartige Kreissektoren zusammen, als Kanten von ihr auslaufen. Alsdann haben wir *drittens* jedesmal an der Ecke $\mathfrak{p}^{(\tau\sigma\varrho)}$ von $\mathfrak{R}$ denjenigen *Sektor* der Kugel mit $\mathfrak{p}^{(\tau\sigma\varrho)}$ als Mittelpunkt vom Radius t zu konstruieren, welcher mit dem Normalensektor dieser Ecke von $\mathfrak{R}$ homothetisch ist; dieser Kugelsektor erhält jene Kreissektoren als begrenzende Seitenflächen. Das Volumen dieses Kugelsektors ist $\frac{1}{3}t^3\chi^{(\tau\sigma\varrho)}$, wenn wir mit $\frac{1}{3}\chi^{(\tau\sigma\varrho)}$ das Volumen des Normalensektors der Ecke $\mathfrak{p}^{(\tau\sigma\varrho)}$ von $\mathfrak{R}$ bezeichnen.

Das Polyeder $\mathfrak{R}$ nun, die hier errichteten Prismen auf den Seitenflächen $\mathfrak{F}^{(\tau)}$, die Zylinderausschnitte an den Kanten $\mathfrak{Q}^{(\tau\sigma)}$, die Kugelsektoren an den Ecken $\mathfrak{p}^{(\tau\sigma\varrho)}$ von $\mathfrak{R}$ stellen lauter solche Bereiche vor, die untereinander in ihren inneren Punkten durchweg verschieden sind, und durch ihre Vereinigung ergeben sie genau den ganzen Körper $\mathfrak{R} + t\mathfrak{G}$. Das Volumen von $\mathfrak{R} + t\mathfrak{G}$ erhält danach den Ausdruck

$$(159) \quad V_0 + t \sum F^{(\tau)} + \frac{1}{2} t^2 \sum \chi^{(\tau\sigma)} L^{(\tau\sigma)} + \frac{1}{3} t^3 \sum \chi^{(\tau\sigma\varrho)},$$

wo die drei Summen beziehungsweise über alle *verschiedenen* Flächen, Kanten, Ecken des Polyeders $\mathfrak{R}$ zu erstrecken sind. Durch Vergleichung mit (158) entsteht danach wieder $3V_1 = \sum F^{(\tau)}$; weiter gewinnen wir, indem wir noch dem Umstande Rechnung tragen, daß jede Kante $\mathfrak{L}^{(\tau\sigma)}$ zwei Seitenflächen $\mathfrak{F}^{(\tau)}$ angehört, für $3V_2$ die Darstellung

$$(160) \quad 3V_2 = \frac{1}{4} \sum_{\tau} \sum_{\sigma} \chi^{(\tau\sigma)} L^{(\tau\sigma)} \quad (\tau = 1, 2, \dots, \nu; \sigma = 1, 2, \dots, \mu_{\tau}).$$

Den hier entwickelten Ausdruck für das Volumen des Raumes zwischen einer polyedrischen Fläche und einer Paralleelfläche zu ihr hat Steiner*) aufgestellt. Nach der Benennung von Steiner würde $6V_2$ als die „Summe der Kantenkrümmung“ von $\mathfrak{R}$ anzusprechen sein. Wir können passender, wie an einer späteren Stelle ausgeführt werden soll, den Wert $\frac{3V_2}{4\pi}$ als den *mittleren Krümmungsradius* von $\mathfrak{R}$ oder, wie wir sogleich näher erklären werden, als den *mittleren Stützebenenabstand* von $\mathfrak{R}$ bezeichnen.

§ 25. Der mittlere Stützebenenabstand eines konvexen Körpers.

Wir wollen uns jetzt zunächst mit der Bildung $V_2 = V(\mathfrak{R}, \mathfrak{G}, \mathfrak{G})$ in bezug auf einen beliebigen konvexen Bezirk $\mathfrak{R}$ beschäftigen und den Wert dieser Größe durch die Stützebenenfunktion von $\mathfrak{R}$ darzustellen suchen.

Zu dem Zwecke haben wir einige Bemerkungen über die Annäherung der Kugel $\mathfrak{G}$ ($x^2 + y^2 + z^2 \leq 1$) durch Polyeder voranzuschicken. Es bedeute $\mathfrak{G}$ die Begrenzung von $\mathfrak{G}$, also die Kugelfläche $x^2 + y^2 + z^2 = 1$. Unter der *Winkeldistanz* zweier Punkte t und t^* auf $\mathfrak{G}$ wollen wir den Winkel (≥ 0 und $\leq \pi$) der zwei Richtungen von o nach t und von o nach t^* verstehen. Ist t ein fester Punkt auf $\mathfrak{G}$ und ϑ irgendein Wert > 0 und $< \pi$, so bildet der Bereich aller Punkte auf $\mathfrak{G}$, die eine Winkeldistanz $\leq \vartheta$ von t haben, eine *Kalotte* von $\mathfrak{G}$, die wir mit $\mathfrak{G}(t, \vartheta)$ bezeichnen, und die sämtlichen Radien von o nach den Punkten dieser Kalotte bilden einen Sektor, den wir $\mathfrak{G}(t, \vartheta)$ nennen. Das Volumen dieses Sektors ist $\frac{2\pi}{3}(1 - \cos \vartheta)$. Für einen Wert $\vartheta < \frac{\pi}{2}$ ist dieser Sektor stets ein konvexer Körper mit o als Eckpunkt, für $\vartheta = \frac{\pi}{2}$ eine Halbkugel von $\mathfrak{G}$.

*) Steiner, Über parallele Flächen, Monatsberichte der Berliner Akademie, 1840, S. 114–118. (Werke, Bd. II, S. 173–176.)

Bedeutet 2ϑ einen beliebigen positiven Wert, den wir $< \frac{\pi}{2}$ annehmen, so können wir auf $\mathfrak{G}$ eine endliche Anzahl von Punkten $t_1, t_2, \dots, t_\nu$ derart bestimmen, daß alsdann jeder beliebige Punkt t der Fläche $\mathfrak{G}$ stets von wenigstens einem dieser ν Punkte eine Winkeldistanz $\leq 2\vartheta$ besitzt oder, wie wir dafür kurz sagen wollen, daß ganz $\mathfrak{G}$ von $t_1, t_2, \dots, t_\nu$ Winkeldistanzen $\leq 2\vartheta$ hat. Diese Forderung ist offenbar gleichbedeutend damit, daß die Kalotten $\mathfrak{G}(t_1, 2\vartheta), \mathfrak{G}(t_2, 2\vartheta), \dots, \mathfrak{G}(t_\nu, 2\vartheta)$ zusammen die ganze Fläche $\mathfrak{G}$ überdecken sollen.

Eine spezielle Methode zur Ermittlung einer solchen Punktreihe $t_1, t_2, \dots, t_\nu$ auf $\mathfrak{G}$ ist die folgende. Wir wählen den ersten Punkt t_1 beliebig auf $\mathfrak{G}$, dann den zweiten Punkt t_2 auf $\mathfrak{G}$ außerhalb der Kalotte $\mathfrak{G}(t_1, 2\vartheta)$, dann t_3 auf $\mathfrak{G}$ außerhalb der Kalotten $\mathfrak{G}(t_1, 2\vartheta)$ und $\mathfrak{G}(t_2, 2\vartheta)$ usw. Haben wir in solcher Art bereits gewisse Punkte $t_1, t_2, \dots, t_{x-1}$ auf $\mathfrak{G}$ angenommen, so wählen wir, wenn die $x-1$ Kalotten $\mathfrak{G}(t_1, 2\vartheta), \mathfrak{G}(t_2, 2\vartheta), \dots, \mathfrak{G}(t_{x-1}, 2\vartheta)$ noch nicht die ganze Fläche $\mathfrak{G}$ überdecken, weiter t_x als einen beliebigen Punkt auf $\mathfrak{G}$ außerhalb jeder dieser $x-1$ Kalotten, also derart, daß alle Winkel

$$t_1 \vartheta t_x, t_2 \vartheta t_x, \dots, t_{x-1} \vartheta t_x > 2\vartheta$$

sind. Die Reihe der Punkte $t_1, t_2, \dots$ läßt sich diesen Forderungen gemäß jedenfalls nicht unbegrenzt fortsetzen. Denn sind wir auf dem angegebenen Wege noch zu t_x gelangt und betrachten wir nun die Kalotten $\mathfrak{G}(t_1, \vartheta), \mathfrak{G}(t_2, \vartheta), \dots, \mathfrak{G}(t_x, \vartheta)$ zu der halb so großen Winkeldistanz ϑ , so haben offenbar irgend zwei dieser letzteren Kalotten niemals einen Punkt miteinander gemein und stoßen daher die dazugehörigen Sektoren $\mathfrak{G}(t_1, \vartheta), \mathfrak{G}(t_2, \vartheta), \dots, \mathfrak{G}(t_x, \vartheta)$ untereinander nur in dem Punkte o zusammen. Diese Sektoren stellen also lauter getrennte Teile der Kugel $\mathfrak{G}$ vor und muß daher die Summe ihrer Volumina kleiner als das Volumen der ganzen Kugel $\mathfrak{G}$ sein, d. h. wir haben

$$(161) \quad \frac{2\pi}{3} x (1 - \cos \vartheta) < \frac{4\pi}{3}$$

und besteht mithin eine obere Grenze für die Zahl x .

Nach dem angegebenen Prinzipie fortschreitend, müssen wir daher schließlich eine Reihe von Punkten $t_1, t_2, \dots, t_\nu$ auf $\mathfrak{G}$ erlangen derart, daß alle Winkel

$$t_i \vartheta t_\nu > 2\vartheta \quad (\iota < x; \iota, x = 1, 2, \dots, \nu)$$

sind, daß aber nunmehr für jeden beliebigen Punkt t auf $\mathfrak{G}$ stets wenigstens einer von den ν Winkeln $tot_1, tot_2, \dots, tot_\nu$ sich $\leq 2\vartheta$ erweist, oder also, daß in der Tat die Kalotten $\mathfrak{G}(t_1, 2\vartheta), \mathfrak{G}(t_2, 2\vartheta), \dots, \mathfrak{G}(t_\nu, 2\vartheta)$ die ganze Kugelfläche $\mathfrak{G}$ überdecken. Dabei wird noch (161) mit $x = \nu$ gelten, andererseits aber erfüllen nun die Sektoren $\mathfrak{G}(t_1, 2\vartheta), \mathfrak{G}(t_2, 2\vartheta), \dots,$

$\mathfrak{G}(t, 2\vartheta)$ die ganze Kugel $\mathfrak{G}$ und wird daher gleichzeitig

$$(162) \quad \frac{2\pi}{3} \nu (1 - \cos 2\vartheta) \geq \frac{4\pi}{3}$$

statthaben. Da wir $2\vartheta < \frac{\pi}{2}$ angenommen haben, können dabei $t_1, t_2, \dots, t_\nu$ nicht sämtlich in einer Ebene durch $\mathfrak{o}$ gelegen sein.

Denken wir uns jetzt eine endliche Reihe von Punkten $t_1, t_2, \dots, t_\nu$ auf $\mathfrak{G}$ irgendwie derart gewählt, daß ganz $\mathfrak{G}$ von ihr Winkeldistanzen $\leq 2\vartheta$ hat. Es seien $\alpha_x, \beta_x, \gamma_x$ die Koordinaten von t_x ($x = 1, 2, \dots, \nu$). Wir bezeichnen mit $\{\mathfrak{T}_x\}$ die Tangentialebene an $\mathfrak{G}$ durch t_x , d. i. die Stützebene an $\mathfrak{G}$ mit der Bedingung

$$\alpha_x x + \beta_x y + \gamma_x z \leq 1;$$

es sei sodann $\mathfrak{S}$ der Bereich, der alle diese ν Ungleichungen für $x = 1, 2, \dots, \nu$ erfüllt. Dieser Bereich $\mathfrak{S}$ enthält die Kugel $\mathfrak{G}$ ganz in sich. Andererseits ist für jeden beliebigen von $\mathfrak{o}$ verschiedenen Punkt $p(x, y, z)$ stets wenigstens einer der ν Winkel $\text{pot}_x \leq 2\vartheta$, somit die zugehörige Größe

$$\alpha_x x + \beta_x y + \gamma_x z \geq \cos 2\vartheta \sqrt{x^2 + y^2 + z^2},$$

und befindet sich danach $\mathfrak{S}$ ganz in der Kugel $\frac{1}{\cos 2\vartheta} \mathfrak{G}$ mit $\mathfrak{o}$ als Mittelpunkt vom Radius $\frac{1}{\cos 2\vartheta}$, ist also ganz im Endlichen gelegen und stellt demnach ein konvexes Polyeder mit den ν extremen Stützebenen $\{\mathfrak{T}_x\}$ vor. Wir bezeichnen mit $\mathfrak{T}_x$ die Seitenfläche von $\mathfrak{S}$ in der Ebene $\{\mathfrak{T}_x\}$, mit T_x den Flächeninhalt dieser Seitenfläche. Die Pyramide $\mathfrak{o}\mathfrak{T}_x$ mit $\mathfrak{o}$ als Spitze und $\mathfrak{T}_x$ als Basis hat mit der Kugel $\mathfrak{G}$ eine gewisse Partie $\mathfrak{G}_x$ gemein, die aus allen den Punkten α, β, γ auf $\mathfrak{G}$ besteht, für welche $\alpha_x \alpha + \beta_x \beta + \gamma_x \gamma$ dem Betrage nach nicht kleiner ist als irgendeiner der $\nu - 1$ anderen Ausdrücke $\alpha_i \alpha + \beta_i \beta + \gamma_i \gamma$ ($i \neq x$), für welche also die Winkeldistanz von t_x nicht größer ist als die Winkeldistanz von irgendeinem der $\nu - 1$ anderen Punkte t_i ($i \neq x$). Mit $\mathfrak{G}$ hat sodann $\mathfrak{o}\mathfrak{T}_x$ ein Gebiet $\mathfrak{G}_x$ gemein, das aus allen Radien von $\mathfrak{o}$ nach den Punkten von $\mathfrak{G}_x$ besteht. Das Volumen dieser Kugelpyramide $\mathfrak{G}_x$, die einen konvexen Körper vorstellt, setzen wir $= \frac{1}{3} \mathfrak{D}_x$ und nennen alsdann $\mathfrak{D}_x$ die Oberfläche von $\mathfrak{G}_x$. Da $\mathfrak{G}_x$ jedenfalls ganz in dem Sektor $\mathfrak{G}(t_x, 2\vartheta)$ enthalten ist, wird

$$\frac{1}{3} \mathfrak{D}_x \leq \frac{2\pi}{3} (1 - \cos 2\vartheta)$$

sein.

Die ν Pyramiden $\mathfrak{o}\mathfrak{T}_x$ ($x = 1, 2, \dots, \nu$) setzen das ganze Polyeder $\mathfrak{S}$ und dementsprechend die ν Kugelpyramiden $\mathfrak{G}_x$ die ganze Kugel $\mathfrak{G}$ zusammen; danach gilt

$$(163) \quad \frac{1}{3} (\mathfrak{D}_1 + \mathfrak{D}_2 + \dots + \mathfrak{D}_v) = \frac{4\pi}{3}.$$

Da $\mathfrak{D}_x$ den Bereich $\mathfrak{G}_x$ enthält und selbst ganz in $\frac{1}{\cos 2\vartheta} \mathfrak{G}_x$ enthalten ist, haben wir

$$(164) \quad \frac{1}{3} \mathfrak{D}_x \leq \frac{1}{3} T_x \leq \frac{1}{3 \cos^2 2\vartheta} \mathfrak{D}_x.$$

Jetzt bedeute $H(\alpha, \beta, \gamma)$ irgendeine Funktion, welche für alle Punkte α, β, γ auf $\mathfrak{G}$ definiert und auf $\mathfrak{G}$ stetig ist. Zu einer beliebig angenommenen positiven Größe ε können wir dann stets einen Winkel $2\vartheta (< \frac{\pi}{2})$ derart bestimmen, daß für je zwei Punkte α, β, γ und $\alpha^*, \beta^*, \gamma^*$ auf $\mathfrak{G}$, deren Winkeldistanz $\leq 2\vartheta$ ist, immer

$$(165) \quad |H(\alpha^*, \beta^*, \gamma^*) - H(\alpha, \beta, \gamma)| \leq \varepsilon$$

wird. Wir ermitteln hierauf irgendeine Punktreihe $t_1, t_2, \dots, t_v$ auf $\mathfrak{G}$, von der ganz $\mathfrak{G}$ Winkeldistanzen $\leq 2\vartheta$ hat, und stellen für sie den Ausdruck

$$(166) \quad H_{\mathfrak{G}} = \frac{1}{4\pi} \sum H(\alpha_x, \beta_x, \gamma_x) \mathfrak{D}_x \quad (x = 1, 2, \dots, v)$$

her, indem wir dabei die Zeichen $\alpha_x, \beta_x, \gamma_x, \mathfrak{D}_x$, wie oben erklärt, verwenden. Ist nun $t_1^*, t_2^*, \dots, t_v^*$ irgendeine zweite Punktreihe auf $\mathfrak{G}$, von der ebenfalls ganz $\mathfrak{G}$ Winkeldistanzen $\leq 2\vartheta$ hat, so entsteht durch Vereinigung sämtlicher verschiedener Punkte aus den beiden Reihen stets wieder eine Punktreihe auf $\mathfrak{G}$ von demselben Charakter. Daraus entnehmen wir leicht mit Rücksicht auf (165) und (163), daß der zu (166) entsprechende Ausdruck für die zweite Punktreihe von dem Ausdrucke (166) selbst subtrahiert eine Differenz ergibt, deren Betrag jedenfalls $\leq 2\varepsilon$ ist. Danach konvergiert der Ausdruck $H_{\mathfrak{G}}$, wenn wir den Winkel ϑ nach Null abnehmen lassen, nach einer bestimmten Grenze, die wir $H_{\mathfrak{G}}$ schreiben. Diese Grenze nennen wir das *Oberflächenintegral* $\frac{1}{4\pi} \int H(\alpha, \beta, \gamma) d\omega$ über die Kugelfläche $\mathfrak{G}$ oder auch den *Mittelwert der Funktion* $H(\alpha, \beta, \gamma)$ für alle möglichen Richtungen α, β, γ , wobei wir in Betracht ziehen, daß für $H(\alpha, \beta, \gamma) = 1$ auch $H_{\mathfrak{G}}$ konstant $= 1$ ist. Von dem Faktor $d\omega$ sprechen wir als dem *Flächenelement* von $\mathfrak{G}$ an der Stelle α, β, γ .

Es bedeute jetzt $\mathfrak{R}$ einen beliebigen konvexen Bezirk und $H(u, v, w)$ dessen Stützebenenfunktion; ferner sei R das Maximum von $H(\alpha, \beta, \gamma)$ für die Punkte α, β, γ auf $\mathfrak{G}$. Für je zwei Punkte $\alpha^*, \beta^*, \gamma^*$ und α, β, γ auf $\mathfrak{G}$ mit einer Winkeldistanz $\leq 2\vartheta$ ist die geradlinige Distanz $\leq 2 \sin \vartheta$ und folgt daher (vgl. die Formel (2) in § 3) stets

$$(167) \quad |H(\alpha^*, \beta^*, \gamma^*) - H(\alpha, \beta, \gamma)| \leq 2R \sin \vartheta.$$

Die Funktion $H(\alpha, \beta, \gamma)$ ist also auf $\mathfrak{E}$ stetig. Den Mittelwert $H_{\mathfrak{G}}$ von $H(\alpha, \beta, \gamma)$ für alle möglichen Richtungen α, β, γ bezeichnen wir als den *mittleren Stützebenenabstand* des konvexen Bezirks $\mathfrak{R}$.

Denken wir uns jetzt zu einem Winkel $2\vartheta < \frac{\pi}{2}$ irgendwie eine Punktreihe $t_1, t_2, \dots, t_r$ auf $\mathfrak{E}$ konstruiert, von der ganz $\mathfrak{E}$ Winkeldistanzen $\leq 2\vartheta$ hat und dazu wie oben das Polyeder $\mathfrak{S}$ gebildet; dabei enthält $\mathfrak{S}$ die Kugel $\mathfrak{G}$ und ist selbst ganz in $\frac{1}{\cos 2\vartheta} \mathfrak{G}$ enthalten. Nach dem Satze in (133) haben wir demzufolge

$$(168) \quad V(\mathfrak{R}, \mathfrak{G}, \mathfrak{G}) \leq V(\mathfrak{R}, \mathfrak{S}, \mathfrak{S}) \leq \frac{1}{\cos^2 2\vartheta} V(\mathfrak{R}, \mathfrak{G}, \mathfrak{G}).$$

Stellen wir nun $V(\mathfrak{S}, \mathfrak{S}, \mathfrak{R})$ gemäß der Formel (135) dar und ziehen die Ungleichungen (164) heran, so erhalten wir unter Verwendung der Abkürzung (166):

$$(169) \quad \frac{4\pi}{3} H_{\mathfrak{S}} \leq V(\mathfrak{S}, \mathfrak{S}, \mathfrak{R}) \leq \frac{4\pi}{3 \cos^3 2\vartheta} H_{\mathfrak{S}}.$$

Die beiden Formeln (168) und (169) zusammen führen dazu, daß für den Betrag der Differenz $V(\mathfrak{R}, \mathfrak{G}, \mathfrak{G}) - \frac{4\pi}{3} H_{\mathfrak{G}}$ eine obere Grenze besteht, deren Größe nach Null konvergiert, wenn wir ϑ nach Null abnehmen lassen. Da nun diese Differenz von ϑ gar nicht abhängt, kommen wir damit zu dem Satze:

Für einen konvexen Bezirk $\mathfrak{R}$ ist $V_2 = V(\mathfrak{R}, \mathfrak{G}, \mathfrak{G})$ gleich dem $\frac{4\pi}{3}$ -fachen des mittleren Stützebenenabstands von $\mathfrak{R}$, also

$$(170) \quad \frac{3}{4\pi} V_2 = \frac{1}{4\pi} \int H(\alpha, \beta, \gamma) d\omega.$$

Vermöge dieser Darstellung erkennen wir insbesondere, daß, wenn ein konvexer Bezirk $\mathfrak{R}$ ganz in einem anderen konvexen Bezirk $\mathfrak{R}^*$ enthalten ist, stets

$$(171) \quad V(\mathfrak{R}, \mathfrak{G}, \mathfrak{G}) < V(\mathfrak{R}^*, \mathfrak{G}, \mathfrak{G})$$

gilt.

Nehmen wir für $\mathfrak{R}$ in (170) beispielsweise eine Strecke $\mathfrak{D}_0$ von der Länge 1, die einen Durchmesser von $\frac{1}{2} \mathfrak{G}$ bildet, also zwei Punkte $-\frac{1}{2} \alpha_0, -\frac{1}{2} \beta_0, -\frac{1}{2} \gamma_0$ und $\frac{1}{2} \alpha_0, \frac{1}{2} \beta_0, \frac{1}{2} \gamma_0$ auf $\frac{1}{2} \mathfrak{E}$ verbindet, so besteht der Körper $\mathfrak{D}_0 + t\mathfrak{G}$ aus dem geraden Kreiszylinder, dessen Achse $\mathfrak{D}_0$ ist und dessen Grundflächen Kreisflächen vom Radius t sind und zwei auf diesen Grundflächen aufgesetzten Halbkugeln. Das Volumen dieses Körpers ist $\pi t^2 + \frac{4\pi}{3} t^3$, also für ihn $3 V_2 = \pi$, andererseits haben wir hier $H(\alpha, \beta, \gamma) = \frac{1}{2} |\alpha \alpha_0 + \beta \beta_0 + \gamma \gamma_0|$ und entsteht demnach

$$(172) \quad \frac{1}{2} = \frac{1}{4\pi} \int |\alpha \alpha_0 + \beta \beta_0 + \gamma \gamma_0| d\omega.$$

Da $\mathfrak{G}$ den Nullpunkt als Mittelpunkt hat, besitzt derjenige konvexe Bezirk $\mathfrak{R}^{(-)}$, der zu einem gegebenen konvexen Bezirk $\mathfrak{R}$ in bezug auf den Nullpunkt symmetrisch ist, die gleiche Invariante V_2 wie $\mathfrak{R}$. Verwandeln wir in (170) α, β, γ und $-\alpha, -\beta, -\gamma$ und addieren die entstehende Formel zu (170), so erhalten wir

$$(173) \quad \frac{3}{4\pi} V_2 = \frac{1}{8\pi} \int (H(\alpha, \beta, \gamma) + H(-\alpha, -\beta, -\gamma)) d\omega.$$

Die Summe $H(\alpha, \beta, \gamma) + H(-\alpha, -\beta, -\gamma)$ hierin bedeutet den gegenseitigen Abstand der zwei parallelen Stützebenen an $\mathfrak{R}$ mit den äußeren Normalen α, β, γ und $-\alpha, -\beta, -\gamma$, die *Breite* von $\mathfrak{R}$ in der Richtung (α, β, γ) . Denken wir uns jetzt den Bezirk $\mathfrak{R}$ als Körper, also alle seine Breiten > 0 . Sind p und q zwei Punkte von $\mathfrak{R}$ in den Stützebenen an $\mathfrak{R}$ mit den Normalen α, β, γ und $-\alpha, -\beta, -\gamma$, so ist die Länge der sie verbindenden Strecke pq mindestens so groß wie die Breite von $\mathfrak{R}$ in der Richtung und gilt daher mit Rücksicht auf das Resultat bei (172) und auf (126):

$$\frac{3}{4\pi} V(pq, \mathfrak{G}, \mathfrak{G}) \geq \frac{1}{4} (H(\alpha, \beta, \gamma) + H(-\alpha, -\beta, -\gamma)).$$

Andererseits ist diese Strecke pq ganz in $\mathfrak{R}$ enthalten und haben wir daher gemäß (171)

$$\frac{3}{4\pi} V(\mathfrak{R}, \mathfrak{G}, \mathfrak{G}) > \frac{3}{4\pi} V(pq, \mathfrak{G}, \mathfrak{G}),$$

so daß stets

$$(174) \quad \frac{3}{4\pi} V_2 > \frac{1}{4} (H(\alpha, \beta, \gamma) + H(-\alpha, -\beta, -\gamma))$$

sein wird. Denken wir uns den Nullpunkt der Koordinaten mit dem Schwerpunkte von $\mathfrak{R}$ zusammenfallend, so gilt nach (143) immer

$$H(-\alpha, -\beta, -\gamma) \geq \frac{1}{3} H(\alpha, \beta, \gamma)$$

und daher

$$(175) \quad \frac{3}{4\pi} V_2 > \frac{1}{3} H(\alpha, \beta, \gamma).$$

Insbesondere wird danach das Maximum von $H(\alpha, \beta, \gamma)$, d. i. der Radius R der kleinsten, den Körper $\mathfrak{R}$ ganz enthaltenden Kugel mit dem Schwerpunkte von $\mathfrak{R}$ als Mittelpunkt stets $< \frac{9}{4\pi} V_2$ sein. Ist $\mathfrak{R}$ speziell ein Körper mit dem Nullpunkt als Mittelpunkt, so besteht allgemein

$$H(-\alpha, -\beta, -\gamma) = H(\alpha, \beta, \gamma)$$

und ist daher

$$R < \frac{3}{2\pi} V_2.$$

Andererseits ist aus (173) unmittelbar ersichtlich, daß die größte vorhandene Breite von $\mathfrak{R}$ sicher $\geq \frac{3}{2\pi} V_2$ ist und hier nur das Gleichheits-

zeichen statthat, wenn alle Breiten von $\mathfrak{K}$ untereinander gleich sind, und aus (170) folgt, daß gewiß $R \geq \frac{3}{4\pi} V_2$ ist und hier das Gleichheitszeichen nur gilt, wenn $\mathfrak{K}$ eine Kugel vorstellt.

§ 26. Die Oberfläche eines konvexen Körpers als das Vierfache des arithmetischen Mittels seiner Projektionen.

Es sei $\mathfrak{K}$ ein konvexer Körper mit der Stützebenenfunktion $H(u, v, w)$. Für jeden Punkt x, y, z aus $\mathfrak{K}$ genügen die Werte x, y den Bedingungen (176)

$$ux + vy \leq H(u, v, 0)$$

für alle möglichen Wertsysteme u, v . Umgekehrt, ist a, b ein spezielles System, wofür stets $H(u, v, 0) - au - bv \geq 0$ ausfällt, so läßt sich nach den Ausführungen in § 10 immer eine Konstante c derart bestimmen, daß auch $H(u, v, w) - au - bv - cw \geq 0$ für beliebige Werte u, v, w gilt, und alsdann ist der Punkt mit den Koordinaten a, b, c ein Punkt aus $\mathfrak{K}$. Den durch die sämtlichen Ungleichungen (176) definierten Bereich $\mathfrak{J}$ in der xy -Ebene, ein gewisses Oval in dieser Ebene, bezeichnen wir danach als die (orthogonale) *Projektion* des Körpers $\mathfrak{K}$ auf die xy -Ebene. Wir wollen jetzt den Flächeninhalt F dieses Ovals $\mathfrak{J}$ als ein gemischtes Volumen darstellen.

Es bedeute $\mathfrak{C}$ die Strecke von der Länge 1, die vom Punkte $0, 0, -\frac{1}{2}$ nach dem Punkte $0, 0, \frac{1}{2}$ führt und betrachten wir den Körper $\mathfrak{K} + t\mathfrak{C}$, wobei wir uns den Parameter $t > H(0, 0, 1) + H(0, 0, -1) = t_0$ denken. Für jeden Punkt a, b, c in $\mathfrak{K}$ ist $-H(0, 0, -1) \leq c \leq H(0, 0, 1)$. Während a, b, c den Körper $\mathfrak{K}$ beschreibt, wird der Körper $\mathfrak{K} + t\mathfrak{C}$ von der Gesamtheit aller daraus abzuleitenden Punkte a, b, z erzeugt, für welche $-\frac{t}{2} \leq z - c \leq \frac{t}{2}$ gilt. Dabei verbleibt nun bei jedem beliebigen Systeme a, b aus $\mathfrak{J}$ die Koordinate z stets in den Grenzen

$$-\frac{t}{2} - H(0, 0, -1) \leq z \leq \frac{t}{2} + H(0, 0, 1),$$

und andererseits wird, zufolge unserer Voraussetzung über t , dabei z jedesmal zum mindesten alle Werte des Intervalls

$$-\frac{t}{2} + H(0, 0, 1) \leq z \leq \frac{t}{2} - H(0, 0, -1)$$

annehmen. Damit haben wir hier zwei gerade Zylinder mit $\mathfrak{J}$ als Querschnitt und Höhen $= t \pm t_0$ erlangt, von denen der erste den Körper $\mathfrak{K} + t\mathfrak{C}$ enthält, der zweite ganz in diesem Körper enthalten ist und finden wir danach das Volumen von $\mathfrak{K} + t\mathfrak{C}$ einerseits $\leq (t + t_0)F$, andererseits $\geq (t - t_0)F$. Stellen wir dieses Resultat mit dem Ausdrucke zusammen, der für dieses Volumen gemäß (134) besteht, und lassen wir t über jede

Grenze hinauswachsen, so erkennen wir, daß $V(\mathfrak{E}, \mathfrak{E}, \mathfrak{E}) = 0$, $3V(\mathfrak{R}, \mathfrak{E}, \mathfrak{E}) = 0$, $3V(\mathfrak{R}, \mathfrak{R}, \mathfrak{E}) = F$ ist.

In ganz analoger Weise können wir die (orthogonale) Projektion des Körpers $\mathfrak{R}$ auf eine Ebene $\alpha x + \beta y + \gamma z = 0$ mit einer beliebigen Richtung (α, β, γ) als Normale definieren. Wir brauchen nur eine orthogonale Koordinatentransformation heranzuziehen, wobei die neue dritte Achse in jene Richtung (α, β, γ) fällt. Die betreffende Projektion, die ein Oval in jener Ebene wird, nennen wir $\mathfrak{F}(\alpha, \beta, \gamma)$ und ihren Flächeninhalt $F(\alpha, \beta, \gamma)$; wir finden alsdann, wenn r den Punkt mit den Koordinaten α, β, γ und (or) die Strecke von o nach r bedeutet, jedesmal

$$3V(\mathfrak{R}, \mathfrak{R}, (or)) = F(\alpha, \beta, \gamma).$$

Wir definieren nun eine Funktion $F(u, v, w)$ für alle möglichen reellen Werte von u, v, w , indem wir, wenn $t > 0$ ist, allgemein

$$(178) \quad F(t\alpha, t\beta, t\gamma) = tF(\alpha, \beta, \gamma)$$

festsetzen und zudem noch $F(0, 0, 0) = 0$ nehmen. Bedeutet $\mathfrak{f}$ den Punkt mit den Koordinaten u, v, w und $(o\mathfrak{f})$ die Strecke von o nach $\mathfrak{f}$, bzw. wenn $\mathfrak{f}$ mit o identisch ist, diesen Punkt allein, so haben wir alsdann stets

$$(179) \quad 3V(\mathfrak{R}, \mathfrak{R}, (o\mathfrak{f})) = F(u, v, w).$$

Auf Grund dieser Formel können wir jetzt allgemein die Ungleichung

$$(180) \quad F(u_1, v_1, w_1) + F(u_2, v_2, w_2) \geq F(u_1 + u_2, v_1 + v_2, w_1 + w_2)$$

erweisen. In der Tat, es seien $\mathfrak{f}_1, \mathfrak{f}_2, \mathfrak{f}$ die Punkte mit den Koordinaten u_1, v_1, w_1 ; u_2, v_2, w_2 und $u_1 + u_2, v_1 + v_2, w_1 + w_2$. Der im Sinne von § 17 aus $(o\mathfrak{f}_1)$ und $(o\mathfrak{f}_2)$ durch Addition hervorgehende Bereich bedeutet dann das Parallelogramm mit den Ecken $o, \mathfrak{f}_1, \mathfrak{f}_2, \mathfrak{f}$ und enthält daher jedenfalls den Bereich $(o\mathfrak{f})$ ganz in sich; nach (133) ist deshalb

$$3V(\mathfrak{R}, \mathfrak{R}, (o\mathfrak{f}_1)) + (o\mathfrak{f}_2) \geq 3V(\mathfrak{R}, \mathfrak{R}, (o\mathfrak{f})),$$

während die linke Seite hier sich nach (127)

$$= 3V(\mathfrak{R}, \mathfrak{R}, (o\mathfrak{f}_1)) + 3V(\mathfrak{R}, \mathfrak{R}, (o\mathfrak{f}_2))$$

ergibt. Diese Umstände in Verbindung mit (179) führen sofort zur Ungleichung (180).

Ferner gilt, da allgemein je zwei Projektionen $\mathfrak{F}(-\alpha, -\beta, -\gamma)$ und $\mathfrak{F}(\alpha, \beta, \gamma)$ identisch sind, immer $F(-u, -v, -w) = F(u, v, w)$.

Nach ihren hier entwickelten Eigenschaften stellt nun $F(u, v, w)$ die Stützebenenfunktion eines gewissen konvexen Körpers $\mathfrak{L}$ mit dem Nullpunkte als Mittelpunkt dar, den wir den *Körper der Projektionen* von $\mathfrak{R}$ nennen wollen. Ersetzen wir $\mathfrak{R}$ durch $t\mathfrak{R}$, wobei $t > 0$ ist, so geht $\mathfrak{L}$ in $t^2\mathfrak{L}$ über; unterwerfen wir $\mathfrak{R}$ einer Translation, so bleibt $\mathfrak{L}$ unverändert. Wir werden jetzt zeigen, daß die Oberfläche des Körpers $\mathfrak{L}$ gleich dem Vierfachen des mittleren Stützebenenabstandes von $\mathfrak{L}$ ist.

Betrachten wir zunächst den Fall, daß $\mathfrak{K}$ ein Polyeder ist; es besitze $\mathfrak{K}$ dabei ν Seitenflächen $\mathfrak{F}_\tau (\tau = 1, 2, \dots, \nu)$ und es sei $(\alpha_\tau, \beta_\tau, \gamma_\tau)$ die äußere Normale und F_τ der Flächeninhalt von $\mathfrak{F}_\tau$. Die Oberfläche von $\mathfrak{K}$ wird definiert durch

$$(181) \quad \mathfrak{O} = 3 V(\mathfrak{K}, \mathfrak{K}, \mathfrak{G}) = F_1 + F_2 + \dots + F_\nu.$$

Fassen wir die orthogonale Projektion von $\mathfrak{K}$ auf irgendeine Ebene $\alpha x + \beta y + \gamma z = 0$ ins Auge, so trifft eine auf dieser Ebene senkrechte Gerade durch einen beliebigen inwendigen Punkt dieser Projektion $\mathfrak{F}(\alpha, \beta, \gamma)$ die Begrenzung von $\mathfrak{K}$ in zwei Punkten, einmal in einem Punkte einer solchen Seitenfläche $\mathfrak{F}_\tau$, für die $\alpha\alpha_\tau + \beta\beta_\tau + \gamma\gamma_\tau > 0$ ausfällt, und dann in einem Punkte einer solchen Seitenfläche $\mathfrak{F}_\tau$, für die der entsprechende Ausdruck < 0 ausfällt, und infolgedessen ergibt sich der doppelte Flächeninhalt jener Projektion

$$(182) \quad 2F(\alpha, \beta, \gamma) = \sum_{\tau} |\alpha\alpha_\tau + \beta\beta_\tau + \gamma\gamma_\tau| F_\tau \quad (\tau = 1, 2, \dots, \nu).$$

Mit Benutzung von (172) finden wir daraus den mittleren Stützebenenabstand des Körpers $\mathfrak{L}$:

$$\frac{1}{4\pi} \int F(\alpha, \beta, \gamma) d\omega = \frac{1}{4} \sum_{\tau} F_\tau = \frac{1}{4} \mathfrak{O},$$

also die Relation

$$(183) \quad \frac{3}{4\pi} V(\mathfrak{L}, \mathfrak{G}, \mathfrak{G}) = \frac{3}{4} V(\mathfrak{K}, \mathfrak{K}, \mathfrak{G}).$$

Diese Beziehung können wir jetzt leicht auf einen beliebigen konvexen Körper $\mathfrak{K}$ ausdehnen, indem wir folgende Bemerkung verwenden: Es seien $\mathfrak{K}, \mathfrak{K}^*$ zwei konvexe Körper, $\mathfrak{L}, \mathfrak{L}^*$ die Körper der Projektionen von $\mathfrak{K}$, bzw. von $\mathfrak{K}^*$ und H, H^*, F, F^* die Stützebenenfunktionen von $\mathfrak{K}, \mathfrak{K}^*, \mathfrak{L}, \mathfrak{L}^*$. Ist der Körper $\mathfrak{K}$ ganz im Körper $\mathfrak{K}^*$ enthalten, so leuchtet ein, daß auch die Projektion von $\mathfrak{K}$ auf irgendeine Ebene $\alpha x + \beta y + \gamma z = 0$ stets ganz in der Projektion von $\mathfrak{K}^*$ auf diese Ebene enthalten sein und infolgedessen stets $F(\alpha, \beta, \gamma) \leq F^*(\alpha, \beta, \gamma)$ und mithin auch $\mathfrak{L}$ in $\mathfrak{L}^*$ enthalten sein wird. Wir können noch hinzufügen, daß, wenn dabei $\mathfrak{K}$ und $\mathfrak{K}^*$ verschieden sind, auch $\mathfrak{L}$ und $\mathfrak{L}^*$ nicht identisch sein können. Denn es läßt sich dann gewiß eine solche Richtung $(\alpha_0, \beta_0, \gamma_0)$ angeben, für die $H(\alpha_0, \beta_0, \gamma_0) < H^*(\alpha_0, \beta_0, \gamma_0)$ ist. Wählen wir nun für (α, β, γ) irgendeine zu $(\alpha_0, \beta_0, \gamma_0)$ orthogonale Richtung, so erweisen sich die orthogonalen Projektionen von $\mathfrak{K}$ und von $\mathfrak{K}^*$ auf die Ebene $\alpha x + \beta y + \gamma z = 0$ in dieser als zwei Ovale mit verschiedenen Stützgeradenfunktionen und muß deshalb $F(\alpha, \beta, \gamma) < F^*(\alpha, \beta, \gamma)$ sein.

Ist jetzt $\mathfrak{K}$ ein beliebiger konvexer Körper und kein Polyeder, so greifen wir einen inneren Punkt $\mathfrak{k}$ von $\mathfrak{K}$ beliebig heraus und können dann zu jeder positiven Größe ε ein Polyeder $\mathfrak{K}^*$ derart bestimmen, daß $\mathfrak{K}$

ganz in $\mathfrak{R}^*$ enthalten ist und andererseits das Polyeder $\frac{1}{1+\varepsilon}\mathfrak{R}^* + \frac{\varepsilon}{1+\varepsilon}\mathfrak{I}$ ganz enthält. Bezeichnen wir mit $\mathfrak{L}^*$ den Körper der Projektionen von $\mathfrak{R}^*$, so ist dann jedesmal $\mathfrak{L}$ in $\mathfrak{L}^*$ enthalten und enthält andererseits den Körper $\frac{1}{(1+\varepsilon)^2}\mathfrak{L}^*$. Wir haben mithin

$$\frac{3}{4\pi} V(\mathfrak{L}^*, \mathfrak{G}, \mathfrak{G}) \geq \frac{3}{4\pi} V(\mathfrak{L}, \mathfrak{G}, \mathfrak{G}) \geq \frac{3}{4(1+\varepsilon)^2\pi} V(\mathfrak{L}^*, \mathfrak{G}, \mathfrak{G}),$$

während zugleich

$$\frac{3}{4} V(\mathfrak{R}^*, \mathfrak{R}^*, \mathfrak{G}) \geq \frac{3}{4} V(\mathfrak{R}, \mathfrak{R}, \mathfrak{G}) \geq \frac{3}{4(1+\varepsilon)^2} V(\mathfrak{R}^*, \mathfrak{R}^*, \mathfrak{G})$$

sein wird. Nun wissen wir bereits, da $\mathfrak{R}^*$ ein Polyeder ist, daß in diesen beiden Reihen bzw. die an erster und die an dritter Stelle aufgeführten Glieder miteinander übereinstimmen. Indem wir die Größe ε nach Null abnehmen lassen, erhalten wir alsdann die Formel (183) auch für den Körper $\mathfrak{R}$. Die damit allgemein festgestellte Formel

$$(184) \quad \frac{1}{4} \mathfrak{D} = \frac{1}{4\pi} \int F(\alpha, \beta, \gamma) d\omega$$

sprechen wir kurz in folgender Weise aus:

Die Oberfläche eines konvexen Körpers ist gleich dem Vierfachen des arithmetischen Mittels der Querschnitte aller dem Körper umbeschriebenen Zylinder.

Aus dieser Darstellung der Oberfläche eines konvexen Körpers sehen wir mit Rücksicht auf die bei (171) gemachte Bemerkung insbesondere sofort, daß, wenn ein konvexer Körper $\mathfrak{R}$ in einem anderen konvexen Körper $\mathfrak{R}^*$ enthalten ist, stets die Oberfläche von $\mathfrak{R}$ wirklich kleiner als die Oberfläche von $\mathfrak{R}^*$ ist.

Bringen wir die am Schlusse von § 25 für Körper mit Mittelpunkt gefundenen Ungleichungen $\frac{3}{2\pi} V_2 > R \geq \frac{3}{4\pi} V_2$ auf den oben betrachteten Körper $\mathfrak{L}$ in Anwendung, so ergibt sich, daß das Maximum von $F(\alpha, \beta, \gamma)$ noch $< \frac{1}{2} \mathfrak{D}$ und andererseits $\geq \frac{1}{4} \mathfrak{D}$ ist. Für einen konvexen Körper $\mathfrak{R}$ mit der Oberfläche $\mathfrak{D}$ besitzt danach jede Projektion einen Flächeninhalt $< \frac{1}{2} \mathfrak{D}$ und gibt es stets solche Projektionen, bei welchen der Flächeninhalt $\geq \frac{1}{4} \mathfrak{D}$ ist.

Durch die Relation (172) finden wir überhaupt für eine beliebige ebene Fläche $\mathfrak{F}$ von einem Flächeninhalt F das (wie in (184) hergestellte) arithmetische Mittel ihrer Projektionen auf die sämtlichen Ebenen durch den Nullpunkt jedesmal $= \frac{1}{2} F$. Infolgedessen können wir auch bei beliebigen Flächen im Raume, die weder konvex noch geschlossen zu sein

brauchen, unter sehr allgemeinen Bedingungen die Oberfläche gleich dem Doppelten des arithmetischen Mittels aller orthogonaler Projektionen der Fläche setzen, wobei wir aber bei jeder Projektion die mehrfach überdeckten Partien nach der Zahl der Überdeckungen in Rechnung zu bringen haben.

§ 27. Verallgemeinerung des Oberflächenbegriffs.

Während die letzten Betrachtungen ausschließlich auf den speziellen Eigenschaften der Kugel beruhten, entwickeln wir jetzt noch eine bemerkenswerte Verallgemeinerung des Begriffs der Oberfläche, wobei die Rolle, welche die Kugel $\mathfrak{G}$ in der Formel $\mathfrak{D} = 3 V(\mathfrak{R}, \mathfrak{R}, \mathfrak{G})$ des § 24 spielt, durch einen beliebigen konvexen Bezirk übernommen wird.

Wir denken uns zu jeder Richtung (α, β, γ) , also für jeden Punkt der Kugelfläche $\alpha^2 + \beta^2 + \gamma^2 = 1$, einen Wert $M(\alpha, \beta, \gamma)$ vorläufig in völlig beliebiger Weise angenommen; dieser Wert kann positiv, Null oder negativ sein. Bedeutet dann $\mathfrak{F}$ ein Polygon in einer Ebene $\alpha x + \beta y + \gamma z = d$ mit der Normale (α, β, γ) und ist F der Flächeninhalt dieses Polygons im gewöhnlichen Sinne, so wollen wir als *den verallgemeinerten* (kurz geschrieben: v.) *Flächeninhalt von $\mathfrak{F}$ auf der (α, β, γ) -Seite jener Ebene* das Produkt $M(\alpha, \beta, \gamma) F$ definieren. Dabei werden, wenn die Faktoren $M(\alpha, \beta, \gamma)$ und $M(-\alpha, -\beta, -\gamma)$ verschieden gewählt sind, dem Polygone $\mathfrak{F}$ verschiedene v. Flächeninhalte auf den zwei Seiten seiner Ebene zugeschrieben. Bei einem Polyeder $\mathfrak{P}$ bezeichnen wir als die *äußeren* v. Flächeninhalte der Seitenflächen deren v. Flächeninhalte auf denjenigen Seiten ihrer Ebenen, welche die äußeren Normalen dieser Ebenen in bezug auf $\mathfrak{P}$ vorstellen.

Wir definieren nun als verallgemeinerte (v.) Oberfläche eines Polyeders die Summe der äußeren v. Flächeninhalte seiner sämtlichen Seitenflächen.

Nunmehr stellen wir die Frage, welche Bedingungen der Funktion $M(\alpha, \beta, \gamma)$ für die Gesamtheit aller Richtungen (α, β, γ) aufzuerlegen sind, damit auch jedem beliebigen konvexen Körper eine bestimmte Größe der v. Oberfläche zugewiesen werden kann und dabei dann allgemein der folgenden Forderung Genüge geschieht:

(I.) *Wenn ein konvexer Körper $\mathfrak{R}$ in einem anderen konvexen Körper $\mathfrak{R}^*$ enthalten ist, soll die v. Oberfläche von $\mathfrak{R}$ niemals größer als die v. Oberfläche von $\mathfrak{R}^*$ sein.*

Um die angeregte Frage zu erledigen, fassen wir zunächst einige einfache konvexe Körper ins Auge. Es bedeute (α, β, γ) irgendeine Richtung, (or) die Strecke vom Nullpunkte $\mathfrak{o}$ nach dem Punkte $\mathfrak{r}$ mit den Koordinaten α, β, γ , und $\mathfrak{F}$ ein Dreieck in der Ebene $\alpha x + \beta y + \gamma z = 0$

vom Flächeninhalte $F = 1$ und noch mit dem Nullpunkte o im Inneren. Der Bereich $s\mathfrak{F} + (or)$ bei positivem s bedeutet dann das dreiseitige Prisma mit $s\mathfrak{F}$ als Basis und (or) als Höhe, und dieser konvexe Körper wird um so umfassender, je größer wir s wählen. Die v. Oberfläche dieses Prismas besitzt offenbar einen Ausdruck

$$(M(\alpha, \beta, \gamma) + M(-\alpha, -\beta, -\gamma))s^2 + Ns,$$

worin N einen von s unabhängigen Wert vorstellt, und dieser Ausdruck darf nun nach unserer Forderung mit wachsendem s niemals abnehmen. Daraus geht hervor, daß bei jeder beliebigen Richtung (α, β, γ) notwendig immer

$$(185) \quad M(\alpha, \beta, \gamma) + M(-\alpha, -\beta, -\gamma) \geq 0$$

sein muß.

Wir definieren jetzt eine Funktion $M(u, v, w)$ für beliebige Werte u, v, w , indem wir festsetzen, daß, wenn $t > 0$ ist, stets

$$(186) \quad M(t\alpha, t\beta, t\gamma) = tM(\alpha, \beta, \gamma)$$

sei, und noch $M(0, 0, 0) = 0$ annehmen.

Es seien jetzt $\alpha_1, \beta_1, \gamma_1; \alpha_2, \beta_2, \gamma_2$ irgend zwei Richtungen, die verschieden und auch nicht einander entgegengesetzt sind, t_1, t_2 Werte > 0 und wir setzen

$$t_1\alpha_1 = u_1, \quad t_1\beta_1 = v_1, \quad t_1\gamma_1 = w_1; \quad t_2\alpha_2 = u_2, \quad t_2\beta_2 = v_2, \quad t_2\gamma_2 = w_2.$$

Sodann bedeute $\alpha_3, \beta_3, \gamma_3$ eine der zwei auf jenen beiden senkrechten Richtungen, r_3 den Punkt mit den Koordinaten $\alpha_3, \beta_3, \gamma_3$, und schreiben wir $\psi_i = \alpha_i x + \beta_i y + \gamma_i z$ ($i = 1, 2, 3$). In der Ebene $\psi_3 = 0$ bestimmen wir einen Punkt p_1 auf der Geraden $\psi_1 = 0$, so daß für ihn $\psi_2 < 0$ und die Länge $op_1 = t_1$ ist, und einen Punkt p_2 auf der Geraden $\psi_2 = 0$, so daß für ihn $\psi_1 < 0$ und die Länge $op_2 = t_2$ ist. Im Dreieck $p_1 p_2$ haben dann die Seiten op_1 und op_2 die Längen t_1 und t_2 und die äußeren Normalen $(\alpha_1, \beta_1, \gamma_1)$ und $(\alpha_2, \beta_2, \gamma_2)$; es sei sodann t die Länge der Seite $p_1 p_2$ und $(-\alpha, -\beta, -\gamma)$ deren äußere Normale in bezug auf das Dreieck, und setzen wir $t\alpha = u, t\beta = v, t\gamma = w$. Bringen wir die drei Relationen (136) auf das dreiseitige Prisma mit $p_1 p_2$ als der einen Grundfläche und or_3 als Höhe in Anwendung, so kommen wir, da die zwei Grundflächen des Prismas gleichen Flächeninhalt und entgegengesetzt gerichtete Normalen haben, zu den Gleichungen:

$$u_1 + u_2 - u = 0, \quad v_1 + v_2 - v = 0, \quad w_1 + w_2 - w = 0.$$

Jetzt sei q irgendein Punkt in $\psi_3 = 0$ im Winkelraume $\psi_1 < 0, \psi_2 < 0$ und noch außerhalb des Dreiecks $p_1 p_2$ gelegen. Vergleichen wir die v. Oberflächen für zwei gerade Prismen miteinander, von denen das eine das Viereck $op_1 q p_2$, das andere das Dreieck $p_1 q p_2$ als Basis hat, während die Höhe für beide dieselbe und an Länge $= s$ sein soll, so ent-

hält das erste dieser Prismen das zweite in sich und muß deshalb auf Grund unserer Forderung die Differenz dieser zwei v. Oberflächen ≥ 0 sein. Diese Differenz erhält den Ausdruck

$$s(t_1 M(\alpha_1, \beta_1, \gamma_1) + t_2 M(\alpha_2, \beta_2, \gamma_2) - t M(\alpha, \beta, \gamma)) \\ + F(M(\alpha_3, \beta_3, \gamma_3) + M(-\alpha_3, -\beta_3, -\gamma_3)),$$

wo F den Flächeninhalt des Dreiecks $p_1 p_2 p_3$ bedeutet. Indem nun diese Differenz bei beliebigem positivem Werte s stets ≥ 0 sein soll, muß der Koeffizient von s darin ≥ 0 , also

$$(187) \quad M(u_1, v_1, w_1) + M(u_2, v_2, w_2) \geq M(u_1 + u_2, v_1 + v_2, w_1 + w_2)$$

sein. — Für solche Argumente $u_1, v_1, w_1; u_2, v_2, w_2$, wobei $\frac{u_2}{u_1} = \frac{v_2}{v_1} = \frac{w_2}{w_1}$ und dieser Wert > 0 bzw. < 0 ist, erhalten wir die nämliche Relation aus (186) allein, bzw. aus (186) unter Berücksichtigung von (185).

Für das Bestehen der Forderung (I) muß demnach die Funktion $M(u, v, w)$ den Bedingungen (185), (186), (187) Genüge leisten; sie erweist sich damit nach § 11 als die Stützebenenfunktion irgendeines konvexen Bezirks $\mathfrak{M}$, der einen konvexen Körper, aber auch ein Oval oder eine Strecke oder einen Punkt vorstellen kann. Gemäß der Formel (135) bedeutet dann für ein Polyeder $\mathfrak{P}$ die v. Oberfläche einfach den Wert des Ausdrucks $3V(\mathfrak{P}, \mathfrak{P}, \mathfrak{M})$, und durch den Satz bei (133) erkennen wir, daß alsdann für einen beliebigen konvexen Körper $\mathfrak{R}$ die v. Oberfläche nur mit $3V(\mathfrak{R}, \mathfrak{R}, \mathfrak{M})$ zusammenfallen kann, wenn (I) gelten soll. Definieren wir aber allgemein für einen konvexen Körper $\mathfrak{R}$ die v. Oberfläche durch $3V(\mathfrak{R}, \mathfrak{R}, \mathfrak{M})$, so leuchtet in der Tat ein, daß alle an diesen Begriff hier gestellten Forderungen erfüllt sind. Den gerade verwandten Bezirk $\mathfrak{M}$ bezeichnen wir als den *Eichbezirk* der v. Oberflächen.

Nehmen wir z. B. für $\mathfrak{M}$ die Strecke von o nach dem Punkte c mit den Koordinaten $0, 0, 1$, so bedeutet nach § 26 die v. Oberfläche $3V(\mathfrak{R}, \mathfrak{R}, (oc))$ den Flächeninhalt der Projektion von $\mathfrak{R}$ auf die xy -Ebene.

Nehmen wir ferner für $\mathfrak{M}$ den Würfel $\mathfrak{W}$, der durch

$$(188) \quad 0 \leq x \leq 1, \quad 0 \leq y \leq 1, \quad 0 \leq z \leq 1$$

bestimmt wird. Bedeuten a, b, c die Punkte mit den Koordinaten $1, 0, 0; 0, 1, 0; 0, 0, 1$, so ist $\mathfrak{W}$ identisch mit der Summe der drei Strecken $(oa) + (ob) + (oc)$ und auf Grund der Regel (127) finden wir daher die v. Oberfläche $3V(\mathfrak{R}, \mathfrak{R}, \mathfrak{W})$ gleich der Summe der Flächeninhalte der drei Projektionen des Körpers $\mathfrak{R}$ auf die drei Koordinatenebenen.

Nehmen wir endlich für $\mathfrak{M}$ das Oktaeder, welches durch

$$(189) \quad |x| + |y| + |z| \leq 1$$

definiert ist, so bedeutet $M(\alpha, \beta, \gamma)$, das Maximum von $\alpha x + \beta y + \gamma z$ in diesem Oktaeder, den größten Wert unter den drei Beträgen $|\alpha|, |\beta|, |\gamma|$.

Im Hinblick auf die Formel (135) können wir dann $3V(\mathfrak{R}, \mathfrak{R}, \mathfrak{M})$ als die maximale Projektion von $\mathfrak{R}$ auf die Gesamtheit der drei Koordinatenebenen bezeichnen.

Wir untersuchen jetzt noch, welche besondere Eigentümlichkeit dem Eichbezirk $\mathfrak{M}$ der verallgemeinerten Oberflächen zukommen muß, damit an Stelle von (I) die schärfere Aussage treten kann: Wenn ein konvexer Körper $\mathfrak{R}$ in einem anderen konvexen Körper $\mathfrak{R}^*$ enthalten ist, so erweist sich $3V(\mathfrak{R}, \mathfrak{R}, \mathfrak{M}) < (\text{nicht bloß } \leq) 3V(\mathfrak{R}^*, \mathfrak{R}^*, \mathfrak{M})$. Wir werden für diesen weiteren Umstand als notwendig und hinreichend erkennen, daß der Bezirk $\mathfrak{M}$ keinen Eckpunkt (s. § 13) besitze.

In der Tat, nehmen wir zunächst an, daß $\mathfrak{M}$ irgendeinen Eckpunkt p aufweise. Es sei (α, β, γ) die Normale einer Eckstützebene durch p an $\mathfrak{M}$, so können wir jedenfalls drei Stützebenen durch p an $\mathfrak{M}$ mit drei unabhängigen Normalenrichtungen $(\alpha_i, \beta_i, \gamma_i)$ ($i = 1, 2, 3$), wobei also die Determinante dieser drei Reihen $\alpha_i, \beta_i, \gamma_i$ von Null verschieden ausfällt, derart finden, daß

(190) $\alpha = t_1\alpha_1 + t_2\alpha_2 + t_3\alpha_3, \quad \beta = t_1\beta_1 + t_2\beta_2 + t_3\beta_3, \quad \gamma = t_1\gamma_1 + t_2\gamma_2 + t_3\gamma_3$
mit drei positiven Faktoren t_1, t_2, t_3 gilt. Indem t_1, t_2, t_3 positiv sind, besteht nach (187) und (186) jedenfalls die Ungleichung

$$(191) \quad M(\alpha, \beta, \gamma) \leq t_1 M(\alpha_1, \beta_1, \gamma_1) + t_2 M(\alpha_2, \beta_2, \gamma_2) + t_3 M(\alpha_3, \beta_3, \gamma_3)$$

und da alle jene vier Stützebenen durch den Punkt p gehen, so wird hierin notwendig das Zeichen = gelten.

Setzen wir nun

$$\psi_i = \alpha_i x + \beta_i y + \gamma_i z, \quad \psi = t_1\psi_1 + t_2\psi_2 + t_3\psi_3 \quad (i = 1, 2, 3)$$

und es seien $\mathfrak{f}_1, \mathfrak{f}_2, \mathfrak{f}_3$ die drei Punkte in der Ebene $\psi = -1$, für welche $\psi_2 = 0, \psi_3 = 0$ bzw. $\psi_1 = 0, \psi_3 = 0$ bzw. $\psi_1 = 0, \psi_2 = 0$ gilt. Die vier Punkte $\mathfrak{o}, \mathfrak{f}_1, \mathfrak{f}_2, \mathfrak{f}_3$ sind dann die Eckpunkte eines Tetraeders $\mathfrak{T}$, dessen vier Seitenflächen in die Ebenen $\psi = -1, \psi_1 = 0, \psi_2 = 0, \psi_3 = 0$ fallen und die äußeren Normalen $(-\alpha, -\beta, -\gamma), (\alpha_1, \beta_1, \gamma_1), (\alpha_2, \beta_2, \gamma_2), (\alpha_3, \beta_3, \gamma_3)$ haben. Bezeichnen wir die Flächeninhalte dieser Seitenflächen mit F, F_1, F_2, F_3 , so bestehen nach (136) die Gleichungen

$$F\alpha = F_1\alpha_1 + F_2\alpha_2 + F_3\alpha_3, \quad F\beta = F_1\beta_1 + F_2\beta_2 + F_3\beta_3, \\ F\gamma = F_1\gamma_1 + F_2\gamma_2 + F_3\gamma_3;$$

da die Determinante der rechts stehenden Ausdrücke in F_1, F_2, F_3 von Null verschieden ist, führt der Vergleich mit (190) zu:

$$(192) \quad F_1 = t_1 F, \quad F_2 = t_2 F, \quad F_3 = t_3 F.$$

Ist jetzt $\mathfrak{f}^*$ irgendein Punkt, für den $\psi_1 < 0, \psi_2 < 0, \psi_3 < 0, \psi < -1$ ist, so setzen die zwei Tetraeder $\mathfrak{o}\mathfrak{f}_1\mathfrak{f}_2\mathfrak{f}_3$ und $\mathfrak{f}^*\mathfrak{f}_1\mathfrak{f}_2\mathfrak{f}_3$ sich wieder zu einem konvexen Körper $\mathfrak{T}^*$ zusammen und haben wir für

den Ausdruck

$$3V(\mathfrak{X}^*, \mathfrak{X}^*, \mathfrak{M}) - 3V(\mathfrak{X}, \mathfrak{X}, \mathfrak{M})$$

$$(193) \quad F_1 M(\alpha_1, \beta_1, \gamma_1) + F_2 M(\alpha_2, \beta_2, \gamma_2) \\ + F_3 M(\alpha_3, \beta_3, \gamma_3) - FM(\alpha, \beta, \gamma),$$

welcher vermöge (192) und, da in (191) das Zeichen $=$ gilt, sich $= 0$ erweist. Danach würden hier also in $\mathfrak{X}$ und $\mathfrak{X}^*$ zwei verschiedene konvexe Körper vorliegen, von denen der erste ganz im zweiten enthalten ist und welche doch die gleiche v. Oberfläche besitzen.

Jetzt nehmen wir andererseits an, daß der Bezirk $\mathfrak{M}$ keinen Eckpunkt besitze. Betrachten wir dann irgendein Tetraeder $\mathfrak{X}$, bei welchem die vier Seitenflächen die Flächeninhalte F, F_1, F_2, F_3 und die äußeren Normalen $(-\alpha, -\beta, -\gamma), (\alpha_1, \beta_1, \gamma_1), (\alpha_2, \beta_2, \gamma_2), (\alpha_3, \beta_3, \gamma_3)$ haben mögen, so muß der mit diesen Größen gebildete Ausdruck (193) jedesmal notwendig > 0 sein. Denn es gelten dann die Gleichungen (190) mit (192) und besteht daher (191), ist also der Ausdruck (193) sicher ≥ 0 . Würde nun dieser Ausdruck $= 0$ ausfallen, so würde die Gleichung

$$M(\alpha, \beta, \gamma) - \alpha x - \beta y - \gamma z = 0$$

in $\mathfrak{M}$, da hier durchweg

$$M(\alpha_i, \beta_i, \gamma_i) - \alpha_i x - \beta_i y - \gamma_i z \geq 0 \quad (i=1, 2, 3)$$

ist, nur so eintreten können, daß zugleich in diesen drei letzten Ungleichungen immer das Zeichen $=$ statthat; d. h. der Schnittpunkt der drei Stützebenen an $\mathfrak{M}$ mit den Normalen $(\alpha_i, \beta_i, \gamma_i)$ müßte notwendig ein Punkt von $\mathfrak{M}$ sein und dieser Punkt wäre dann ein Eckpunkt von $\mathfrak{M}$.

Jetzt seien $\mathfrak{R}$ und $\mathfrak{R}^*$ irgend zwei verschiedene konvexe Körper und es sei $\mathfrak{R}$ ganz in $\mathfrak{R}^*$ enthalten. Wir können alsdann jedenfalls eine solche Ebene $\{\mathfrak{F}\}$ konstruieren, welche innere Punkte von $\mathfrak{R}^*$ enthält, aber $\mathfrak{R}$ gar nicht trifft, und es sei (α, β, γ) die Normale dieser Ebene nach der Seite, auf der $\mathfrak{R}$ nicht liegt. Wir nehmen in $\{\mathfrak{F}\}$ drei nicht in gerader Linie gelegene Punkte q_1, q_2, q_3 aus $\mathfrak{R}^*$ an und können hernach einen vierten Punkt q aus $\mathfrak{R}^*$ so nahe an der Ebene $\{\mathfrak{F}\}$ auf ihrer (α, β, γ) -Seite wählen, daß in dem Tetraeder $\mathfrak{X} = (q, q_1, q_2, q_3)$ die äußeren Normalen der drei, q enthaltenden Seitenflächen $\mathfrak{F}_1, \mathfrak{F}_2, \mathfrak{F}_3$ beliebig wenig von (α, β, γ) abweichen und die Ebenen $\{\mathfrak{F}_1\}, \{\mathfrak{F}_2\}, \{\mathfrak{F}_3\}$ dieser Seitenflächen $\mathfrak{R}$ ebenfalls nicht treffen und auf derselben Seite wie jedesmal die vierte Ecke von $\mathfrak{X}$ liegen lassen. Es sei nun $\mathfrak{Z}^*$ derjenige Teil von $\mathfrak{R}^*$, welcher zugleich die Bedingungen der drei Stützebenen $\{\mathfrak{F}_1\}, \{\mathfrak{F}_2\}, \{\mathfrak{F}_3\}$ an $\mathfrak{X}$ erfüllt, und $\mathfrak{Z}$ der Teil von $\mathfrak{Z}^*$, der nicht auf der (α, β, γ) -Seite von $\{\mathfrak{F}\}$ liegt. Alsdann ist $\mathfrak{Z}^*$ in $\mathfrak{R}^*$ und $\mathfrak{R}$ in $\mathfrak{Z}$ enthalten und $\mathfrak{Z}^*$ setzt sich aus $\mathfrak{Z}$ und $\mathfrak{X}$ zusammen, wobei diese Körper in der Fläche $\mathfrak{F}$

aneinanderstoßen. Wir haben daher

$$3V(\mathfrak{R}^*, \mathfrak{R}^*, \mathfrak{M}) \geq 3V(\mathfrak{S}^*, \mathfrak{S}^*, \mathfrak{M}), \quad 3V(\mathfrak{S}, \mathfrak{S}, \mathfrak{M}) \geq 3V(\mathfrak{R}, \mathfrak{R}, \mathfrak{M})$$

und die Differenz $3V(\mathfrak{S}^*, \mathfrak{S}^*, \mathfrak{M}) - 3V(\mathfrak{S}, \mathfrak{S}, \mathfrak{M})$ erweist sich nach (137) gleich dem Ausdrucke (193) in bezug auf das Tetraeder $\mathfrak{T}$, also > 0 , mithin gilt auch stets

$$(194) \quad 3V(\mathfrak{R}^*, \mathfrak{R}^*, \mathfrak{M}) > 3V(\mathfrak{R}, \mathfrak{R}, \mathfrak{M}).$$

Insbesondere erhalten wir auf diese Weise für die Oberflächen im gewöhnlichen Sinne, deren Eichkörper die Kugel vom Radius 1 mit dem Nullpunkt als Mittelpunkt ist, von neuem die Tatsache:

Wenn ein konvexer Körper $\mathfrak{R}$ ganz in einem anderen konvexen Körper $\mathfrak{R}^*$ enthalten ist, so besitzt $\mathfrak{R}$ stets eine kleinere Oberfläche als $\mathfrak{R}^*$.

§ 28. Tangentialkörper und Kappenkörper.

Wir besprechen hier für eine spätere Anwendung noch einen weiteren Grenzfall der Ungleichung (133). Es seien $\mathfrak{R}$ und $\mathfrak{R}'$ zwei konvexe Körper mit den Stützebenenfunktionen H und H' und es sei $\mathfrak{R}'$ ganz in $\mathfrak{R}$ enthalten. Für die vier Größen

$$V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}), \quad V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}'), \quad V(\mathfrak{R}, \mathfrak{R}', \mathfrak{R}'), \quad V(\mathfrak{R}', \mathfrak{R}', \mathfrak{R}'),$$

die wir mit V_0, V_1, V_2, V_3 bezeichnen, bestehen dann nach (133) jedenfalls die Ungleichungen

$$(195) \quad V_0 \geq V_1 \geq V_2 \geq V_3.$$

Ist $\mathfrak{R}'$ nicht identisch mit $\mathfrak{R}$, so gibt es innere Punkte von $\mathfrak{R}$, die nicht zu $\mathfrak{R}'$ gehören und ist infolgedessen gewiß $V_0 > V_3$. Wir fragen jetzt, unter welchen Umständen $V_0 > V_1$ und $V_0 > V_2$ gilt.

Wir wollen einen konvexen Körper $\mathfrak{R}$ als einen *Tangentialkörper* eines konvexen Körpers $\mathfrak{R}'$ bezeichnen, wenn jede Tangentialebene (extreme Stützebene) an $\mathfrak{R}$ stets auch eine Stützebene an $\mathfrak{R}'$ ist. Dabei versteht es sich nach § 14 und § 13 ohne weiteres, daß der Körper $\mathfrak{R}$ den Körper $\mathfrak{R}'$ ganz in sich enthält. Wir werden jetzt den Satz beweisen:

Ist $\mathfrak{R}'$ in $\mathfrak{R}$ enthalten, so besteht dann und nur dann die Gleichung

$$(196) \quad V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) = V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}'),$$

wenn $\mathfrak{R}$ ein Tangentialkörper von $\mathfrak{R}'$ ist.

In der Tat, nehmen wir zunächst an, während $\mathfrak{R}'$ in $\mathfrak{R}$ enthalten ist, gäbe es irgendeine solche Tangentialebene an $\mathfrak{R}$, die *nicht* zugleich Stützebene an $\mathfrak{R}'$ ist. Der Einfachheit wegen denken wir uns den Nullpunkt ins Innere von $\mathfrak{R}'$ und die positive z -Achse in die Richtung der äußeren Normale jener Tangentialebene gelegt. Es sei c_0' der kleinste, c_1' der größte Wert von z in $\mathfrak{R}'$ und c_0 der kleinste, c_1 der größte Wert von z

in $\mathfrak{R}$. Nach Voraussetzung ist dann $c_0' < 0 < c_1' < c_1$. Bedeutet δ eine positive Größe $< c_1 - c_1'$ und bezeichnen wir den Teil von $\mathfrak{R}$, wo $z \leq c_1 - \delta$ ist, mit $\mathfrak{R}_0$, den Teil von $\mathfrak{R}$, wo $z \geq c_1 - \delta$ ist, mit $\mathfrak{R}_1$, so ist $\mathfrak{R}'$ ganz in $\mathfrak{R}_0$ enthalten. Die Ebene $z = c_1 - \delta$ hat mit $\mathfrak{R}$ ein gewisses Oval $\mathfrak{F}(\delta)$ gemein; wir bezeichnen mit $F(\delta)$ den Flächeninhalt, mit $U(\delta)$ den Umfang von $\mathfrak{F}(\delta)$, mit $r(\delta)$ das Maximum unter den Radien aller ganz im Oval $\mathfrak{F}(\delta)$ enthaltenen Kreisflächen. Daß die Stützebene $z = c_1$ an $\mathfrak{R}$ eine Tangentialebene an $\mathfrak{R}$ ist, bedeutet nach § 14 soviel, wie daß der Quotient $\frac{r(\delta)}{\delta}$ für ein nach Null abnehmendes δ über jede Grenze hinauswächst.

Indem das Oval $\mathfrak{F}(\delta)$ eine Kreisfläche vom Radius $r(\delta)$ enthält, ist der gemischte Flächeninhalt von $\mathfrak{F}(\delta)$ und dieser Kreisfläche gewiß nicht größer als der Flächeninhalt von $\mathfrak{F}(\delta)$, d. h. hat man

$$\frac{1}{2} r(\delta) U(\delta) \leq F(\delta)$$

und wird mithin $\frac{F(\delta)}{\delta U(\delta)}$ mit nach Null abnehmendem δ gleichfalls über jede Grenze hinauswachsen.

Der Beziehung (137) gemäß gilt

$$(197) \quad V(\mathfrak{R}_0, \mathfrak{R}_0, \mathfrak{R}') + V(\mathfrak{R}_1, \mathfrak{R}_1, \mathfrak{R}') = V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}') + \frac{1}{3} (c_1' - c_0') F(\delta).$$

Ist $\mathfrak{R}$ und damit auch $\mathfrak{R}_0$ ein Polyeder und stellen wir $V(\mathfrak{R}_0, \mathfrak{R}_0, \mathfrak{R})$ und $V(\mathfrak{R}_0, \mathfrak{R}_0, \mathfrak{R}')$ gemäß (135) dar, so erkennen wir die Differenz der ersten und der zweiten Größe als Aggregat von lauter Termen ≥ 0 , unter denen ein Term $= \frac{1}{3} F(\delta) (c_1 - c_1')$ ist, und wir erhalten demnach

$$(198) \quad V(\mathfrak{R}_0, \mathfrak{R}_0, \mathfrak{R}) \geq V(\mathfrak{R}_0, \mathfrak{R}_0, \mathfrak{R}') + \frac{1}{3} (c_1 - c_1') F(\delta).$$

Diese Formel überträgt sich dann genau wie die Beziehung (137) von dem Falle eines Polyeders auf den Fall eines beliebigen konvexen Körpers $\mathfrak{R}$.

Nach (133) haben wir ferner, da $\mathfrak{R}_0$ in $\mathfrak{R}$ enthalten ist

$$(199) \quad V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) \geq V(\mathfrak{R}_0, \mathfrak{R}_0, \mathfrak{R}).$$

Endlich brauchen wir eine gewisse obere Grenze für $V(\mathfrak{R}_1, \mathfrak{R}_1, \mathfrak{R}')$. Es sei p ein Punkt aus $\mathfrak{R}$ in der Stützebene $z = c_1$, q der Schnittpunkt von op mit $z = c_1 - \delta$, und s der Sinus des Neigungswinkels von op gegen diese Ebene, also $\frac{\delta}{s}$ die Länge von pq . Wir bilden durch Projektion des Ovals $\mathfrak{F}(\delta)$ vom Punkte p aus auf die Ebene $z = c_1$ in dieser das Oval $\frac{c_1}{c_1 - \delta} \mathfrak{F}(\delta)$; konstruieren wir sodann den Zylinder $\mathfrak{C}_1$ mit letzterem Oval als einer Grundfläche und mit Erzeugenden des Mantels, die parallel und gleich pq sind, so daß die andere Grundfläche durch Dilatation des Ovals $\mathfrak{F}(\delta)$ von q aus im Verhältnisse $c_1 : c_1 - \delta$ entsteht,

so enthält dieser Zylinder $\mathfrak{C}_1$ offenbar den Körper $\mathfrak{R}_1$ ganz in sich und gilt daher

$$(200) \quad V(\mathfrak{C}_1, \mathfrak{C}_1, \mathfrak{R}') \geq V(\mathfrak{R}_1, \mathfrak{R}_1, \mathfrak{R}').$$

Denken wir uns nun zunächst $\mathfrak{R}$ als Polyeder, so werden auch $\mathfrak{R}_1$ und $\mathfrak{C}_1$ Polyeder und hat der Mantel von $\mathfrak{C}_1$ einen Flächeninhalt

$$\leq \frac{\delta}{s} \frac{c_1}{c_1 - \delta} U(\delta),$$

wo $U(\delta)$ den Umfang von $\mathfrak{F}(\delta)$ bedeutet. Bringen wir nun zur Berechnung von $V(\mathfrak{C}_1, \mathfrak{C}_1, \mathfrak{R}')$ die Formel (135) in Anwendung und bezeichnen das Maximum der Stützebenenabstände $H'(\alpha, \beta, \gamma)$ für $\mathfrak{R}'$ mit R' , so gelangen wir zur Ungleichung

$$(201) \quad \frac{1}{3} \left(\frac{c_1}{c_1 - \delta} \right)^2 F(\delta)(c_1' - c_0') + \frac{1}{3} \frac{\delta}{s} \frac{c_1}{c_1 - \delta} U(\delta) R' \geq V(\mathfrak{C}_1, \mathfrak{C}_1, \mathfrak{R}').$$

Diese Formel läßt sich dann ebenso wie (137) sofort auf den Fall eines beliebigen konvexen Körpers $\mathfrak{R}$ ausdehnen.

Durch Summation der Beziehungen (197)–(201) entsteht nun

$$(202) \quad V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) - V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}') \geq \frac{1}{3} F(\delta) \left[c_1 - c_1' - \frac{\delta(2c_1 - \delta)}{(c_1 - \delta)^2} (c_1' - c_0') \right] - \frac{1}{3} \delta U(\delta) \frac{c_1 R'}{s(c_1 - \delta)}.$$

Hierin ist $c_1 - c_1' > 0$ und, wie oben bemerkt, konvergiert $\frac{\delta U(\delta)}{F(\delta)}$ mit δ zugleich nach Null. Danach fällt der in (202) rechts vom Zeichen $\geq$ befindliche Ausdruck sicher > 0 aus, wenn δ unter eine gewisse positive Größe gesunken ist, und ergibt sich somit bei der gegenwärtigen Voraussetzung:

$$(203) \quad V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) > V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}').$$

Hiernach ist für die Gleichheit der Werte $V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R})$ und $V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}')$, wenn $\mathfrak{R}'$ in $\mathfrak{R}$ enthalten ist, jedenfalls eine notwendige Bedingung, daß eine jede Tangentialebene von $\mathfrak{R}$ stets zugleich eine Stützebene an $\mathfrak{R}'$ sei.

Nehmen wir jetzt andererseits diese Bedingung als erfüllt, somit $\mathfrak{R}$ als Tangentialkörper von $\mathfrak{R}'$ an. Stellt $\mathfrak{R}$ zunächst ein Polyeder vor, so gilt für die Richtungen (α, β, γ) , welche die äußeren Normalen der Seitenflächen von $\mathfrak{R}$ abgeben, immer $H(\alpha, \beta, \gamma) = H'(\alpha, \beta, \gamma)$, und im Hinblick auf (135) erhalten wir daher in der Tat $V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) = V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}')$. Wenn aber $\mathfrak{R}$ einen beliebigen konvexen Körper bedeutet und wie wir annahmen, $\mathfrak{o}$ ein innerer Punkt von $\mathfrak{R}$ ist, können wir nach dem am Schlusse von § 13 bewiesenen Satze in bezug auf jede vorgegebene Größe ε stets ein solches Polyeder $\mathfrak{P}$ konstruieren, dessen Seitenflächen lauter Tangentialebenen an $\mathfrak{R}$ sind und welches dabei ganz in $(1 + \varepsilon)\mathfrak{R}$ liegt. Da alsdann $\mathfrak{P}$ gleichfalls einen Tangentialkörper von $\mathfrak{R}'$ vorstellt, gilt nach dem soeben

Bemerkten gewiß $V(\mathfrak{P}, \mathfrak{P}, \mathfrak{P}) = V(\mathfrak{P}, \mathfrak{P}, \mathfrak{R}')$. Nun sinken die Differenzen $V(\mathfrak{P}, \mathfrak{P}, \mathfrak{P}) - V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R})$ und $V(\mathfrak{P}, \mathfrak{P}, \mathfrak{R}') - V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}')$, die jedenfalls ≥ 0 sind, mit nach Null abnehmendem ε unter jede positive GröÙe, und danach finden wir auch $V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) - V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}')$ beliebig klein, d. h. gleich Null. —

Wir bezeichnen einen konvexen Körper $\mathfrak{R}$ als einen *Kappenkörper* eines konvexen Körpers $\mathfrak{R}'$, wenn, von den Eckstützebenen an $\mathfrak{R}$ abgesehen, alle übrigen Stützebenen an $\mathfrak{R}$ zugleich auch Stützebenen an $\mathfrak{R}'$ sind (s. § 16). Dabei ist $\mathfrak{R}$ jedenfalls auch ein Tangentialkörper von $\mathfrak{R}'$ und enthält $\mathfrak{R}'$ ganz in sich. Wir werden jetzt den ferner Satz beweisen:

Ist $\mathfrak{R}'$ in $\mathfrak{R}$ enthalten, so besteht dann und nur dann die Gleichung

$$(204) \quad V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) = V(\mathfrak{R}, \mathfrak{R}', \mathfrak{R}'),$$

wenn $\mathfrak{R}$ ein Kappenkörper von $\mathfrak{R}'$ ist.

In der Tat, nehmen wir an, es sei $\mathfrak{R}'$ nicht identisch mit $\mathfrak{R}$, und betrachten wir beliebige zwei Stützebenen an $\mathfrak{R}$ und $\mathfrak{R}'$ mit derselben äußeren Normale, welche nicht zusammenfallen. Der Einfachheit wegen denken wir uns o ins Innere von $\mathfrak{R}'$ und die positive z -Achse in die Richtung jener Normale gelegt, und wir wollen in bezug auf $\mathfrak{R}$ und $\mathfrak{R}'$ und diese Richtung die Zeichen $c_0', c_1', c_0, c_1, \delta, \mathfrak{R}_0, \mathfrak{R}_1, \mathfrak{F}(\delta), F(\delta)$ genau wie vorhin verwenden. Indem wir die positive GröÙe $\delta < c_1 - c_1'$ einrichten, wird $\mathfrak{R}'$ ganz in $\mathfrak{R}_0$ enthalten sein und daher jedenfalls

$$(205) \quad V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) \geq V(\mathfrak{R}, \mathfrak{R}_0, \mathfrak{R}_0) \geq V(\mathfrak{R}, \mathfrak{R}', \mathfrak{R}')$$

statthaben. Ferner besteht gemäß (197) die Gleichung

$$(206) \quad V(\mathfrak{R}, \mathfrak{R}_0, \mathfrak{R}_0) + V(\mathfrak{R}, \mathfrak{R}_1, \mathfrak{R}_1) = V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) + \frac{1}{3} (c_1 - c_0) F(\delta).$$

Denken wir uns zunächst $\mathfrak{R}$ und damit auch $\mathfrak{R}_1$ als Polyeder, so stimmt die Stützebenenfunktion von $\mathfrak{R}_1$ für jede Richtung (α, β, γ) , welche (äußere) Normale einer Seitenfläche von $\mathfrak{R}_1$ vorstellt, mit Ausnahme der Normale $(0, 0, -1)$ von $\mathfrak{F}(\delta)$, genau mit dem Werte $H(\alpha, \beta, \gamma)$ für $\mathfrak{R}$ überein, während für die eine Richtung $(0, 0, -1)$ die erstere Funktion $-(c_1 - \delta)$, dagegen $H(0, 0, -1) = -c_0$ ist. Bilden wir nun $V(\mathfrak{R}_1, \mathfrak{R}_1, \mathfrak{R})$ und $V(\mathfrak{R}_1, \mathfrak{R}_1, \mathfrak{R}_1)$ gemäß der Formel (135), so entsteht daher die Beziehung

$$(207) \quad \frac{1}{3} (c_1 - \delta - c_0) F(\delta) = V(\mathfrak{R}, \mathfrak{R}_1, \mathfrak{R}_1) - V(\mathfrak{R}_1, \mathfrak{R}_1, \mathfrak{R}_1).$$

Diese Gleichung läßt sich dann ebenso wie (137) auf den Fall eines beliebigen konvexen Körpers $\mathfrak{R}$ ausdehnen.

Soll nun $V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) = V(\mathfrak{R}, \mathfrak{R}', \mathfrak{R}')$, mithin gemäß (205) auch $= V(\mathfrak{R}, \mathfrak{R}_0, \mathfrak{R}_0)$ sein, so folgt aus (206) und (207), daß

$$(208) \quad V(\mathfrak{R}_1, \mathfrak{R}_1, \mathfrak{R}_1) = \frac{1}{3} \delta F(\delta)$$

statthaben muß. Ist nun p irgendein Punkt aus $\mathfrak{R}$ in seiner Stützebene $z = c_1$, so bedeutet $\frac{1}{3} \delta F(\delta)$ das Volumen des Kegels mit $\mathfrak{F}(\delta)$ als Grundfläche und p als Spitze. Dieser Kegel ist jedenfalls ganz in $\mathfrak{R}_1$ enthalten. Die Relation (208) kann danach nur bestehen, wenn $\mathfrak{R}_1$ mit diesem Kegel identisch wird; alsdann stellt p einen Eckpunkt von $\mathfrak{R}$ und $z = c_1$ eine Eckstützebene durch p an $\mathfrak{R}$ vor.

Auf diese Weise erkennen wir in der Tat als eine notwendige Bedingung für das Eintreten der Gleichung $V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R}) = V(\mathfrak{R}, \mathfrak{R}', \mathfrak{R}')$, während $\mathfrak{R}'$ in $\mathfrak{R}$ enthalten ist, daß jede solche Stützebene an $\mathfrak{R}$, welche nicht Eckstützebene an $\mathfrak{R}$ ist, immer zugleich eine Stützebene an $\mathfrak{R}'$ sein muß, d. h. daß $\mathfrak{R}$ ein Kappenkörper von $\mathfrak{R}'$ ist.

Setzen wir andererseits jetzt $\mathfrak{R}$ als Kappenkörper von $\mathfrak{R}'$ voraus. Stellt $\mathfrak{R}'$ ein Polyeder vor, so ist alsdann, wie in § 16 gezeigt wurde, $\mathfrak{R}$ notwendig ebenfalls ein Polyeder und kommt der vorausgesetzte Charakter von $\mathfrak{R}$ darauf hinaus, daß nicht nur eine jede Seitenfläche von $\mathfrak{R}$, sondern auch bereits eine jede Kante von $\mathfrak{R}$ stets wenigstens einen Punkt von $\mathfrak{R}'$ enthält. Berechnen wir nun $V(\mathfrak{R}, \mathfrak{R}', \mathfrak{R}')$ gemäß (124), indem wir an die Stelle der Polyeder $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ dort jetzt $\mathfrak{R}, \mathfrak{R}', \mathfrak{R}'$ treten lassen, so wird darin ein von Null verschiedener Faktor $F_{12}^{(\tau)}$ jedenfalls nur solchen Richtungen $(\alpha^{(\tau)}, \beta^{(\tau)}, \gamma^{(\tau)})$ entsprechen, wobei die Stützebenen an $\mathfrak{R}$ mit diesen Normalen sei es eine Seitenfläche, sei es eine Kante von $\mathfrak{R}$ enthalten. Für diese Richtungen stimmen nun gewiß die Funktionen H von $\mathfrak{R}$ und H' von $\mathfrak{R}'$ stets überein, und daher folgt $V(\mathfrak{R}, \mathfrak{R}', \mathfrak{R}') = V(\mathfrak{R}, \mathfrak{R}', \mathfrak{R})$ und da $\mathfrak{R}$ zugleich ein Tangentialkörper von $\mathfrak{R}'$ ist, nach (196) dann weiter $= V(\mathfrak{R}, \mathfrak{R}, \mathfrak{R})$. Ist $\mathfrak{R}'$ ein beliebiger konvexer Körper und denken wir uns $\mathfrak{o}$ im Inneren von $\mathfrak{R}'$ gelegen, so können wir in bezug auf jede positive Größe ε zu $\mathfrak{R}'$ stets ein Polyeder $\mathfrak{P}'$ konstruieren, welches $\mathfrak{R}'$ enthält und selbst in $(1 + \varepsilon)\mathfrak{R}'$ enthalten ist. Es sei $\mathfrak{R}$ nicht identisch mit $\mathfrak{R}'$, so liegt bei hinreichend kleinem ε gewiß irgendein Eckpunkt von $\mathfrak{R}$ außerhalb $\mathfrak{P}'$. Bezeichnet p einen Eckpunkt von $\mathfrak{R}$ außerhalb $\mathfrak{P}'$, so ist der Normalensektor des Eckpunktes p von $(p, \mathfrak{P}')$ ganz enthalten im Normalensektor des Eckpunktes p von $(p, \mathfrak{R}')$. Es seien nun $p_1, p_2, \dots$ die sämtlichen Eckpunkte von $\mathfrak{R}$ außerhalb $\mathfrak{P}'$, so werden hiernach, da $\mathfrak{R}$ ein Kappenkörper von $\mathfrak{R}'$ ist, auch die Normalensektoren dieser Punkte in bezug auf $\mathfrak{P}'$ untereinander in den inneren Punkten durchweg verschieden ausfallen und wird daher der gemäß § 16 zu bildende konvexe Körper $\mathfrak{P} = (p_1, p_2, \dots, \mathfrak{P}')$ einen Kappenkörper von $\mathfrak{P}'$ vorstellen; dabei erweist sich die Zahl der Eckpunkte $p_1, p_2, \dots$ jedenfalls als eine endliche und

wird $\mathfrak{P}$ wieder ein Polyeder sein. Dieses Polyeder $\mathfrak{P}$ enthält den Körper $\mathfrak{K}$ und ist seinerseits ganz in $(1 + \varepsilon)\mathfrak{K}$ enthalten. Nach dem vorhin Bewiesenen finden wir bereits

$$V(\mathfrak{P}, \mathfrak{P}, \mathfrak{P}) - V(\mathfrak{P}, \mathfrak{P}', \mathfrak{P}') = 0,$$

und von der Differenz hier links weicht $V(\mathfrak{K}, \mathfrak{K}, \mathfrak{K}) - V(\mathfrak{K}, \mathfrak{K}', \mathfrak{K}')$ um eine Größe ab, für welche eine, zugleich mit ε nach Null konvergierende obere Grenze besteht. Danach muß nun in jedem Falle

$$V(\mathfrak{K}, \mathfrak{K}, \mathfrak{K}) = V(\mathfrak{K}, \mathfrak{K}', \mathfrak{K}')$$

gelten, wenn $\mathfrak{K}$ ein Kappenkörper von $\mathfrak{K}'$ ist.

XXVI.

Volumen und Oberfläche.

Herrn Rudolf Lipschitz zum fünfzigjährigen Doktorjubiläum, 9. August 1903,
in herzlicher Verehrung gewidmet vom Verfasser.
(Mathematische Annalen, Band 57, S. 447—495).

Für die konvexen Körper gibt es einen elementaren Weg, um den Begriff der Oberfläche aus dem einfacheren Begriffe des Volumens heraus zu entwickeln, und in Verfolg dieses Weges gelangt man zu sehr bemerkenswerten Erweiterungen der Tatsache, wonach unter allen Körpern gleichen Volumens die Kugel die kleinste Oberfläche besitzt.

Liegt ein konvexer Körper $\mathfrak{K}$ vor und versteht man unter x, y, z rechtwinklige Koordinaten eines Punktes aus $\mathfrak{K}$, so nimmt ein linearer Ausdruck $ux + vy + wz$, wo u, v, w feste Größen sind, in $\mathfrak{K}$ immer einen bestimmten größten Wert $H(u, v, w)$ an; und diese Funktion $H(u, v, w)$ von drei beliebigen reellen Argumenten, die Stützebenenfunktion von $\mathfrak{K}$, charakterisiert den konvexen Körper $\mathfrak{K}$ vollkommen. Das Volumen des Körpers $\mathfrak{K}$ erscheint als ein gewisser homogener Ausdruck dritten Grades V_H^3 in den sämtlichen Werten $H(u, v, w)$. Aus diesem Ausdrucke entspringt für drei beliebige konvexe Körper $\mathfrak{K}_1, \mathfrak{K}_2, \mathfrak{K}_3$ eine polare Bildung, ein symbolisches Produkt $V_{H_1} V_{H_2} V_{H_3}$, das gemischte Volumen der drei Körper $\mathfrak{K}_1, \mathfrak{K}_2, \mathfrak{K}_3$. Diese Größe ist invariant bei beliebigen Translationen der einzelnen Körper. Werden zwei der Körper mit einem bestimmten Körper $\mathfrak{K}$, der dritte aber mit einer Kugel vom Radius 1 identifiziert, so ist das dreifache ihres gemischten Volumens die Oberfläche von $\mathfrak{K}$.

Für die gemischten Volumina gilt der wichtige Satz: Für irgend drei Körper vom Volumen 1 wird das gemischte Volumen stets ≥ 1 und nur dann $= 1$, wenn die drei Körper miteinander homothetisch sind. Daß jeder konvexe Körper, der keine Kugel ist, eine größere Oberfläche hat als eine Kugel von demselben Volumen, ist nur ein spezieller Fall dieses Satzes.

Diese fundamentale Ungleichung läßt weiter die folgende Auslegung zu: Man bezeichne in der Mannigfaltigkeit aller möglichen Funktionen $H(u, v, w)$ von drei reellen Argumenten u, v, w eine einzelne Funktion

$H(u, v, w)$ als einen „Punkt“, den Inbegriff der aus zwei Funktionen H_1 und H_2 abzuleitenden Funktionen $(1-t)H_1 + tH_2$ für $0 \leq t \leq 1$ als die H_1 und H_2 verbindende „Strecke“; alsdann besitzt die Gesamtheit der Stützebenenfunktionen H zu allen denjenigen konvexen Körpern, welche ein Volumen ≥ 1 haben, die Eigenschaft, mit irgend zwei Punkten stets die ganze sie verbindende Strecke zu enthalten, stellt also ein „konvexes Gebilde“ in jener Mannigfaltigkeit vor.

Geht man auf die Tangentialebenen des Gebildes ein, so ist sein konvexer Charakter gleichbedeutend mit folgendem Theorem:

Auf der Kugelfläche vom Radius 1 mit dem Nullpunkt als Mittelpunkt denke man sich Masse in einer beliebigen stetigen und durchweg positiven Flächendichtigkeit ausgebreitet, doch so, daß der Schwerpunkt der ganzen Belegung in den Nullpunkt fällt; alsdann existiert eine geschlossene konvexe Fläche, bei welcher an jeder Stelle das Produkt der Krümmungsradien gleich der Flächendichtigkeit an dem Punkte der Kugel mit gleicher Normale ist; und diese Fläche ist völlig bestimmt bis auf eine beliebige Translation, durch die man sie noch variieren kann.

In diesem Theorem erkennt man eine Aussage über eine gewisse quadratische partielle Differentialgleichung zweiter Ordnung, deren Lösbarkeit unter bestimmten Bedingungen hier durch eine eigenartige, wohl noch mancher weiteren Anwendungen fähige Methode sichergestellt wird.

§ 1. Stützebenenfunktion eines konvexen Körpers.

1. Es seien x, y, z rechtwinklige Koordinaten eines Punktes im Raume, und $\mathfrak{M}$ bedeute eine *abgeschlossene* Menge von Punkten x, y, z , die ganz in einer Kugel von *endlichem* Radius enthalten ist, aber nicht völlig in eine einzige Ebene fällt. Sind u, v, w irgendwelche festen Werte, so hat der Ausdruck $ux + vy + wz$ für die Gesamtheit der Punkte x, y, z in $\mathfrak{M}$ ein bestimmtes *Maximum*, das $H(u, v, w)$ heiße.

Diese Funktion $H(u, v, w)$ von drei beliebigen reellen Argumenten erfüllt offenbar folgende Bedingungen (1)–(4):

$$(1) \quad H(0, 0, 0) = 0,$$

$$(2) \quad H(tu, tv, tw) = tH(u, v, w),$$

wenn $t > 0$ ist. Sind u_1, v_1, w_1 und u_2, v_2, w_2 irgend zwei Systeme der Argumente, so gibt es in $\mathfrak{M}$ immer wenigstens einen Punkt x, y, z , wofür

$$(u_1 + u_2)x + (v_1 + v_2)y + (w_1 + w_2)z = H(u_1 + u_2, v_1 + v_2, w_1 + w_2)$$

wird, und da für diesen Punkt sicherlich

$$u_1x + v_1y + w_1z \leq H(u_1, v_1, w_1), \quad u_2x + v_2y + w_2z \leq H(u_2, v_2, w_2)$$

ist, so gilt daher immer:

$$(3) \quad H(u_1 + u_2, v_1 + v_2, w_1 + w_2) \leq H(u_1, v_1, w_1) + H(u_2, v_2, w_2).$$

Das Maximum von $-(ux + vy + wz)$ in $\mathfrak{M}$ ist $H(-u, -v, -w)$, also gilt in $\mathfrak{M}$ stets:

$$-H(-u, -v, -w) \leq ux + vy + wz \leq H(u, v, w).$$

Wenn die Werte $u, v, w \neq 0, 0, 0$ sind, muß daher, da $\mathfrak{M}$ nicht ganz in einer Ebene liegen soll, stets

$$(4) \quad H(u, v, w) + H(-u, -v, -w) > 0$$

sein.

2. Eine Ebene, welche wenigstens einen Punkt der Begrenzung von $\mathfrak{M}$ enthält, aber außer den Punkten, die sie mit $\mathfrak{M}$ gemein hat, $\mathfrak{M}$ ganz auf einer Seite von sich liegen läßt, nennen wir eine *Stützebene an $\mathfrak{M}$* .

Ist $H(u, v, w)$ eine beliebige reelle Funktion von drei reellen Argumenten u, v, w , welche allen den Bedingungen (1)–(4) genügt, so bezeichnen wir den Bereich $\mathfrak{R}$ von Punkten x, y, z , welcher durch die sämtlichen Ungleichungen

$$(5) \quad ux + vy + wz \leq H(u, v, w)$$

für alle möglichen Wertsysteme u, v, w definiert ist, als einen *konvexen Körper*.

Die Funktion H nennen wir die *Stützebenenfunktion* von $\mathfrak{R}$, da aus den Ungleichungen (5) offenbar genau die Stützebenen an $\mathfrak{R}$ zu erkennen sind.

Ist $H(u, v, w)$ wie in 1. aus der Punktmenge $\mathfrak{M}$ hergeleitet, so wird der durch die Ungleichungen (5) definierte Bereich $\mathfrak{R}$ der *kleinste, $\mathfrak{M}$ enthaltende konvexe Körper*, d. h. $\mathfrak{R}$ ist ein notwendiger Bestandteil jedes konvexen Körpers, der $\mathfrak{M}$ ganz in sich aufnimmt.

Ist $\mathfrak{R}^*$ ein zweiter konvexer Körper mit der Stützebenenfunktion $H^*(u, v, w)$, so ist dann und nur dann $\mathfrak{R}$ ganz in $\mathfrak{R}^*$ enthalten, wenn stets

$$H(u, v, w) \leq H^*(u, v, w)$$

ausfällt.

3. Ein konvexer Körper ist andererseits völlig durch die Eigenschaften zu charakterisieren, *erstens*, daß jede Gerade mit ihm sei es eine Strecke, sei es einen Punkt, sei es keinen Punkt gemein hat, *zweitens*, daß zu ihm wenigstens vier nicht in einer Ebene gelegene Punkte gehören.

4. Ist p ein beliebiger Punkt, so verstehen wir unter $\mathfrak{R} + p$ den Körper, der aus $\mathfrak{R}$ durch diejenige *Translation* entsteht, durch welche der Nullpunkt nach p gelangt. Sind a, b, c die Koordinaten von p , so wird die Stützebenenfunktion von $\mathfrak{R} + p$:

$$H(u, v, w) + au + bv + cw.$$

Unterwerfen wir $\mathfrak{R}$ einer *Dilatation* vom Nullpunkte aus nach allen

Richtungen in einem festen positiven Verhältnisse $t:1$, so bezeichnen wir den entstehenden Körper mit $t\mathfrak{R}$; eine Stützebenenfunktion wird $tH(u, v, w)$.

5. Wir bezeichnen mit $\mathfrak{G}$ die Kugel $x^2 + y^2 + z^2 \leq 1$, vom Radius 1 mit dem Nullpunkt o als Mittelpunkt, mit $\mathfrak{E}$ die Kugeloberfläche $x^2 + y^2 + z^2 = 1$, mit α, β, γ die Koordinaten eines beliebigen Punktes auf $\mathfrak{E}$, bzw. die *Richtung* vom Nullpunkte nach diesem Punkte. Infolge der Eigenschaft (2) sind alle Werte der Funktion H bereits durch deren Werte $H(\alpha, \beta, \gamma)$ für die Punkte auf $\mathfrak{E}$ bestimmt. Die Ungleichung

$$\alpha x + \beta y + \gamma z \leq H(\alpha, \beta, \gamma)$$

bezeichnen wir als die *Bedingung der Stützebene an $\mathfrak{R}$ mit der äußeren Normale (α, β, γ)* .

Die Funktion $H(u, v, w)$ ist nach den Eigenschaften (1)–(4) eine *stetige* Funktion ihrer Argumente, und besitzen infolgedessen die Werte $H(\alpha, \beta, \gamma)$ auf $\mathfrak{E}$ ein bestimmtes *Maximum* G . Ist ein Wert $H(\alpha, \beta, \gamma) \leq 0$, so ist nach (4) der zugehörige Wert $H(-\alpha, -\beta, -\gamma)$ positiv und von größerem Betrage; G ist daher jedenfalls > 0 . Mit Hilfe von (3) und (2) gewinnen wir die Ungleichung

$$(6) \quad |H(u - u_0, v - v_0, w - w_0) - H(u_0, v_0, w_0)| \leq G\sqrt{u^2 + v^2 + w^2}.$$

§ 2. Annäherung an einen beliebigen konvexen Körper durch vollkommene Ovaloide.

6. Ist $\varphi = 0$ die Gleichung einer Ebene und der konvexe Körper $\mathfrak{R}$ ganz im Bereiche $\varphi \leq 0$ enthalten, so heißt $\varphi \leq 0$ ein *Halbraum* um $\mathfrak{R}$. Ist ein Halbraum $\varphi \leq 0$ um $\mathfrak{R}$ so beschaffen, daß man *nicht* $\varphi = t_1\varphi_1 + t_2\varphi_2$ setzen kann, so daß $t_1 > 0$, $t_2 > 0$ und $\varphi_1 \leq 0$, $\varphi_2 \leq 0$ zwei *verschiedene* Halbräume um $\mathfrak{R}$ sind, so heißt $\varphi \leq 0$ ein *extremer Halbraum* um $\mathfrak{R}$. Die Ebene $\varphi = 0$ ist dann jedenfalls eine Stützebene an $\mathfrak{R}$ und heißt eine *extreme Stützebene* an $\mathfrak{R}$. Ein konvexer Körper mit einer *endlichen* Anzahl von extremen Stützebenen heißt ein (*konvexes*) *Polyeder*.

7. Unter einem *vollkommenen Ovaloid* wollen wir einen konvexen Körper verstehen, bei dem die *Begrenzung* durch eine *analytische* Gleichung in den rechtwinkligen Koordinaten x, y, z definiert wird und überdies in jedem Punkte eine *bestimmte* und immer nur eine *Berührung erster Ordnung* eingehende *Tangentialebene* besitzt.

8. Ist $\mathfrak{R}$ ein beliebiger konvexer Körper mit dem Nullpunkte als *innerem* Punkt und ε eine beliebige positive Größe, so läßt sich stets ein *vollkommenes Ovaloid* $\mathfrak{Q}$ bestimmen, so daß $\mathfrak{Q}$ den Körper $\mathfrak{R}$ enthält und selbst in $(1 + \varepsilon)\mathfrak{R}$ enthalten ist.

Es sei $H(u, v, w)$ die Stützebenenfunktion von $\mathfrak{R}$. Da der Nullpunkt im Inneren von $\mathfrak{R}$ liegt, ist jede Größe $H(\alpha, \beta, \gamma) > 0$. Es sei nun g

das *Minimum* der Werte $H(\alpha, \beta, \gamma)$ auf der Kugelfläche $\mathfrak{E}$, so ist auch $g > 0$. Wir denken uns den ganzen Raum durch ein *Netz* von lauter gleichen *Würfeln* mit einer Kante δ erfüllt. Es sei $\mathfrak{B}$ der Gesamtbereich aller derjenigen Würfel dieses Netzes, welche überhaupt wenigstens einen Punkt von $\mathfrak{K}$ aufnehmen, und $\mathfrak{P}$ der kleinste, diesen Bereich $\mathfrak{B}$ ganz enthaltende konvexe Körper, so ist $\mathfrak{P}$ ein Polyeder, und für jede Richtung (α, β, γ) ist der Abstand derjenigen Stützebene an $\mathfrak{P}$, welche (α, β, γ) als äußere Normale hat, vom Nullpunkte einerseits $> H(\alpha, \beta, \gamma)$, andererseits

$$\leq H(\alpha, \beta, \gamma) + \delta \sqrt{3} \leq H(\alpha, \beta, \gamma) \left(1 + \frac{\delta \sqrt{3}}{g}\right).$$

Danach enthält $\mathfrak{P}$ den Körper $\mathfrak{K}$ im Inneren und ist selbst ganz in $\left(1 + \frac{\delta \sqrt{3}}{g}\right) \mathfrak{K}$ enthalten.

Das Polyeder $\mathfrak{P}$ besitze n Seitenflächen; da $\mathfrak{P}$ den Nullpunkt im Inneren enthält, können wir die Bedingungen dieser n extremen Stützebenen an $\mathfrak{P}$ in der Form

$$(7) \quad \chi_1 \leq 1, \chi_2 \leq 1, \dots, \chi_n \leq 1$$

schreiben, so daß dabei $\chi_1, \chi_2, \dots, \chi_n$ *homogene* lineare Ausdrücke in x, y, z sind. Es sei nun ω eine beliebige positive GröÙe, die wir $> \lg n$ annehmen, und $\mathfrak{Q}$ der durch die Ungleichung

$$(8) \quad \Omega = e^{\omega} \chi_1 + e^{\omega} \chi_2 + \dots + e^{\omega} \chi_n \leq n e^{\omega}$$

bestimmte Bereich.

Dieser Bereich $\mathfrak{Q}$ enthält jedenfalls das durch die Ungleichungen (7) definierte Polyeder $\mathfrak{P}$ in sich. Andererseits ist $\mathfrak{Q}$ ganz in $\left(1 + \frac{\lg n}{\omega}\right) \mathfrak{P}$ enthalten; denn in jedem Punkte außerhalb des letzteren Polyeders erweist sich stets wenigstens eine der GröÙen $\chi_1, \chi_2, \dots, \chi_n$ als $> 1 + \frac{\lg n}{\omega}$ und die rechte Seite in (8) daher als $> n e^{\omega}$. Nach der Lagenbeziehung von $\mathfrak{P}$ zu $\mathfrak{K}$ wird weiter $\mathfrak{Q}$ den Körper $\mathfrak{K}$ enthalten und selbst in

$$\left(1 + \frac{\delta \sqrt{3}}{g}\right) \left(1 + \frac{\lg n}{\omega}\right) \mathfrak{K}$$

enthalten sein. Wir können nun δ so klein und ω so groß annehmen, daß der hier stehende Faktor von $\mathfrak{K}$ sich $\leq 1 + \varepsilon$ erweist.

Die Begrenzung von $\mathfrak{Q}$ ist die *analytische* Fläche $\Omega = n e^{\omega}$. Wir finden den Ausdruck

$$(9) \quad \frac{\partial \Omega}{\partial x} x + \frac{\partial \Omega}{\partial y} y + \frac{\partial \Omega}{\partial z} z = \omega \chi_1 e^{\omega} \chi_1 + \omega \chi_2 e^{\omega} \chi_2 + \dots + \omega \chi_n e^{\omega} \chi_n,$$

mithin auf der Begrenzung von $\mathfrak{Q}$ stets $\geq \omega e^{\omega} - \frac{n-1}{e}$; denn dort ist in jedem Punkte wenigstens eine der GröÙen $\chi \geq 1$ und andererseits gilt

immer $\xi e^t \geq -\frac{1}{e}$. Mit Berücksichtigung von $e^\omega > n$ ersehen wir hieraus, daß auf der Begrenzung von Ω niemals $\frac{\partial \Omega}{\partial x}, \frac{\partial \Omega}{\partial y}, \frac{\partial \Omega}{\partial z}$ gleichzeitig Null sein können, mithin in jedem Punkte dieser Begrenzung stets eine bestimmte Tangentialebene existiert.

Weiter finden wir, wenn x, y, z als lineare Funktionen eines Parameters t dargestellt werden, immer

$$(10) \quad \frac{d^2 \Omega}{dt^2} = \omega^2 \left(\left(\frac{dx_1}{dt} \right)^2 e^\omega x_1 + \left(\frac{dx_2}{dt} \right)^2 e^\omega x_2 + \dots + \left(\frac{dx_n}{dt} \right)^2 e^\omega x_n \right) > 0.$$

Aus dieser Beziehung folgt, daß auf einer beliebigen geradlinigen Strecke der Ausdruck $\Omega(x, y, z)$ seinen größten Wert immer an wenigstens einem der Endpunkte annimmt, daß mithin Ω mit irgend zwei Punkten stets die ganze sie verbindende Strecke enthält. Andererseits ist aus (10) ersichtlich, daß jede Tangentialebene an Ω mit Ω nur eine Berührung erster Ordnung eingeht. Nach allen diesen Umständen besitzt Ω in der Tat die in unserem Satze verlangten Eigenschaften.

9. Auf Grund dieses Satzes können wir weiter zu einem gegebenen konvexen Körper $\mathfrak{R}$, der den Nullpunkt o im Inneren enthält, mit einer Stützebenenfunktion H , immer eine unendliche Reihe von vollkommenen Ovaloiden $\Omega', \Omega'', \dots$ mit solchen Stützebenenfunktionen $Q', Q'', \dots$ herstellen, daß die Reihe der Quotienten

$$\frac{Q'(\alpha, \beta, \gamma)}{H(\alpha, \beta, \gamma)}, \quad \frac{Q''(\alpha, \beta, \gamma)}{H(\alpha, \beta, \gamma)}, \quad \dots$$

nach der Grenze 1 konvergiert und zwar *gleichmäßig* für alle Systeme α, β, γ auf der ganzen Kugelfläche $\mathfrak{E}$. Trifft der hier bezeichnete Umstand zu, so wollen wir sagen, die Reihe der konvexen Körper $\Omega', \Omega'', \dots$ hat den konvexen Körper $\mathfrak{R}$ als *Grenze*, oder *konvergiert* nach $\mathfrak{R}$.

Ist p ein beliebiger Punkt, so bezeichnen wir weiter den Körper $\mathfrak{R} + p$, der p als inneren Punkt enthält, als *Grenze* der Körper

$$\Omega' + p, \Omega'' + p, \dots$$

§ 3. Volumen eines konvexen Körpers.

10. Jedem konvexen Körper kommt ein bestimmtes Volumen zu, ferner ein bestimmter Schwerpunkt, welcher stets ein innerer Punkt des Körpers ist.

11. Wir führen für die Punkte α, β, γ der Kugelfläche $\mathfrak{E}$ Polarkoordinaten ein und setzen

$$\alpha = \sin \vartheta \cos \psi, \quad \beta = \sin \vartheta \sin \psi, \quad \gamma = \cos \vartheta,$$

wobei wir ϑ und ψ in den Grenzen $0 \leq \vartheta \leq \pi$, $0 \leq \psi \leq 2\pi$ annehmen. Wir schreiben ferner

$$\alpha_1 = \frac{\partial \alpha}{\partial \vartheta} = \cos \vartheta \cos \psi, \quad \beta_1 = \frac{\partial \beta}{\partial \vartheta} = \cos \vartheta \sin \psi, \quad \gamma_1 = \frac{\partial \gamma}{\partial \vartheta} = -\sin \vartheta,$$

$$\alpha_2 = \frac{1}{\sin \vartheta} \frac{\partial \alpha}{\partial \psi} = -\sin \psi, \quad \beta_2 = \frac{1}{\sin \vartheta} \frac{\partial \beta}{\partial \psi} = \cos \psi, \quad \gamma_2 = \frac{1}{\sin \vartheta} \frac{\partial \gamma}{\partial \psi} = 0;$$

dabei ergeben die drei Gleichungen

$$(11) \quad \xi = \alpha_1 x + \beta_1 y + \gamma_1 z, \quad \eta = \alpha_2 x + \beta_2 y + \gamma_2 z, \quad \zeta = \alpha x + \beta y + \gamma z$$

stets eine orthogonale Transformation der Koordinaten x, y, z mit einer Determinante $= +1$.

12. Es sei $\mathfrak{R}$ ein vollkommenes Ovaloid, $\mathfrak{F}$ seine begrenzende Fläche, $H(u, v, w)$ die Stützebenenfunktion von $\mathfrak{R}$. Wir schreiben

$$H(\alpha, \beta, \gamma) = H(\vartheta, \psi) = H.$$

Die Stützebene an $\mathfrak{R}$ mit der äußeren Normale (α, β, γ) hat die Gleichung

$$(12) \quad \xi = H(\vartheta, \psi);$$

sie ist hier zugleich Tangentialebene an $\mathfrak{F}$ und berührt $\mathfrak{F}$ in einem bestimmten Punkte p . Die ganze Fläche $\mathfrak{F}$ erscheint damit punktweise, durch parallele Normalen, auf die Kugelfläche $\mathfrak{E}$ bezogen. Wir wollen nun unter $x, y, z, \alpha, \beta, \gamma, \vartheta, \psi$ speziell die betreffenden Bestimmungsstücke für den Punkt p verstehen und die Veränderungen dieser Größen beim Übergang zu einem anderen Punkte auf $\mathfrak{F}$ durch Vorsetzen von Δ andeuten; ferner sollen $\Delta\xi, \Delta\eta, \Delta\zeta$ die Werte der Ausdrücke (11) bedeuten, wenn darin $\Delta x, \Delta y, \Delta z$ an die Stelle von x, y, z treten. Dann gilt auf $\mathfrak{F}$ in einer gewissen Umgebung von p für $\Delta\xi$ eine Entwicklung nach Potenzen von $\Delta\xi, \Delta\eta$:

$$\Delta\xi = -\frac{1}{2} (P\Delta\xi^2 + 2\Sigma\Delta\xi\Delta\eta + T\Delta\eta^2) + (\Delta\xi, \Delta\eta)_3 + \dots;$$

darin bilden die quadratischen Glieder eine definite negative Form, es ist also $P > 0, PT - \Sigma^2 > 0$. Wir können alsdann, da $PT - \Sigma^2 \neq 0$ ist, in einer gewissen Umgebung von p die Werte $\Delta\xi, \Delta\eta$ durch $\frac{\partial \Delta\xi}{\partial \Delta\xi}$ und $\frac{\partial \Delta\xi}{\partial \Delta\eta}$ ausdrücken, welche letzteren Größen sich sofort mittels $\Delta\alpha, \Delta\beta, \Delta\gamma$ darstellen lassen. Daraus erkennen wir, daß die Koordinaten x, y, z des Punktes p von $\mathfrak{F}$, wo die äußere Normale die Richtung α, β, γ hat, und weiter der zugehörige Wert $H(\alpha, \beta, \gamma)$ analytische Funktionen der Größen α, β, γ auf $\mathfrak{E}$ sind.

Da die Ebene (12) Tangentialebene an $\mathfrak{F}$ ist, haben wir

$$(13) \quad \alpha \frac{\partial x}{\partial \vartheta} + \beta \frac{\partial y}{\partial \vartheta} + \gamma \frac{\partial z}{\partial \vartheta} = 0, \quad \frac{1}{\sin \vartheta} \left(\alpha \frac{\partial x}{\partial \psi} + \beta \frac{\partial y}{\partial \psi} + \gamma \frac{\partial z}{\partial \psi} \right) = 0,$$

und mit Rücksicht hieraus folgen aus den allgemeinen Formeln (11) zur Bestimmung der Koordinaten x, y, z des Punktes p die Gleichungen:

$$(14) \quad \xi = \frac{\partial H(\vartheta, \psi)}{\partial \vartheta}, \quad \eta = \frac{1}{\sin \vartheta} \frac{\partial H(\vartheta, \psi)}{\partial \psi}, \quad \zeta = H(\vartheta, \psi).$$

13. Um nun das Volumen V des Körpers $\mathfrak{K}$ auszudrücken, zerlegen wir die Kugelfläche $\mathfrak{E}$ in Flächenelemente $d\omega = \sin \vartheta d\vartheta d\psi$; jedem Element $d\omega$ entspricht als Abbild durch parallele Normalen ein Element df auf der Fläche $\mathfrak{F}$, und wir konstruieren jedesmal die Pyramide mit dem Nullpunkt $\mathfrak{o}$ als Spitze und dem Element df als Grundfläche; die Höhe dieser Pyramide, mit gewissem Vorzeichen genommen, ist $= H(\alpha, \beta, \gamma)$ und ihr Volumen mit demselben Vorzeichen daher

$$(15) \quad \frac{1}{3} Hdf = \frac{1}{3} \left| x, \frac{\partial x}{\partial \vartheta}, \frac{\partial x}{\partial \psi} \right| d\vartheta d\psi.$$

(Wir bezeichnen hier und weiterhin eine dreireihige Determinante, in welcher die Glieder der ersten Reihe von der Koordinate x abhängen und die der zweiten und dritten in der entsprechenden Weise mit Hilfe des Zeichens y bzw. z darzustellen sind, einfach durch Angabe bloß der ersten Reihe). Der Körper $\mathfrak{K}$ ist nun derart das Aggregat aller jener Elementarpyramiden, daß sein Volumen genau

$$(16) \quad V = \frac{1}{3} \int Hdf = \frac{1}{3} \iint \left| x, \frac{\partial x}{\partial \vartheta}, \frac{\partial x}{\partial \psi} \right| d\vartheta d\psi$$

wird, wo die Integrale über die ganze Fläche $\mathfrak{F}$, bzw. die ganze Kugelfläche $\mathfrak{E}$ zu erstrecken sind.

14. Wir setzen jetzt

$$\alpha \frac{\partial^2 x}{\partial \vartheta^2} + \beta \frac{\partial^2 y}{\partial \vartheta^2} + \gamma \frac{\partial^2 z}{\partial \vartheta^2} = -R, \quad \frac{1}{\sin \vartheta} \left(\alpha \frac{\partial^2 x}{\partial \vartheta \partial \psi} + \beta \frac{\partial^2 y}{\partial \vartheta \partial \psi} + \gamma \frac{\partial^2 z}{\partial \vartheta \partial \psi} \right) = -S,$$

$$\frac{1}{\sin^2 \vartheta} \left(\alpha \frac{\partial^2 x}{\partial \psi^2} + \beta \frac{\partial^2 y}{\partial \psi^2} + \gamma \frac{\partial^2 z}{\partial \psi^2} \right) = -T.$$

Durch Differentiation der zwei Gleichungen (13) einmal nach ϑ , einmal nach ψ erhalten wir noch die Beziehungen

$$R = \alpha_1 \frac{\partial x}{\partial \vartheta} + \beta_1 \frac{\partial y}{\partial \vartheta} + \gamma_1 \frac{\partial z}{\partial \vartheta}, \quad S = \frac{1}{\sin \vartheta} \left(\alpha_1 \frac{\partial x}{\partial \psi} + \beta_1 \frac{\partial y}{\partial \psi} + \gamma_1 \frac{\partial z}{\partial \psi} \right),$$

$$S = \alpha_2 \frac{\partial x}{\partial \vartheta} + \beta_2 \frac{\partial y}{\partial \vartheta} + \gamma_2 \frac{\partial z}{\partial \vartheta}, \quad T = \frac{1}{\sin \vartheta} \left(\alpha_2 \frac{\partial x}{\partial \psi} + \beta_2 \frac{\partial y}{\partial \psi} + \gamma_2 \frac{\partial z}{\partial \psi} \right).$$

Nun gilt

$$\Delta x = \frac{\partial x}{\partial \vartheta} \Delta \vartheta + \frac{\partial x}{\partial \psi} \Delta \psi + \frac{1}{2} \left(\frac{\partial^2 x}{\partial \vartheta^2} \Delta \vartheta^2 + 2 \frac{\partial^2 x}{\partial \vartheta \partial \psi} \Delta \vartheta \Delta \psi + \frac{\partial^2 x}{\partial \psi^2} \Delta \psi^2 \right) + \dots, \dots;$$

aus den Gleichungen (11) gewinnen wir dann mit Rücksicht auf die letzten Ausdrücke und auf (13):

$$\Delta \xi = R \Delta \vartheta + S \sin \vartheta \Delta \psi + \dots, \quad \Delta \eta = S \Delta \vartheta + T \sin \vartheta \Delta \psi + \dots,$$

$$\Delta \zeta = -\frac{1}{2} (R \Delta \vartheta^2 + 2 S \sin \vartheta \Delta \vartheta \Delta \psi + T \sin^2 \vartheta \Delta \psi^2) + \dots,$$

und hieraus geht durch Elimination von $\Delta \vartheta$ und $\Delta \psi$ eine Entwicklung

$$\Delta \xi = -\frac{1}{2} \left(\frac{T \Delta \xi^2 - 2 S \Delta \xi \Delta \eta + R \Delta \eta^2}{R T - S^2} \right) + (\Delta \xi, \Delta \eta)_2 + \dots$$

hervor.

Wir entnehmen daraus für die in 12. benutzten Größen P, Σ, T :

$$P = \frac{T}{RT - S^2}, \quad \Sigma = \frac{-S}{RT - S^2}, \quad T = \frac{R}{RT - S^2},$$

so daß

$$RT - S^2 = \frac{1}{PT - \Sigma^2}$$

das Produkt,

$$R + T = \frac{P + T}{PT - \Sigma^2}$$

die Summe der Hauptkrümmungsradien der Fläche $\mathfrak{F}$ im Punkte p darstellen, während von $\frac{2S}{R-T}$ die Neigung der Krümmungskurven auf $\mathfrak{F}$ durch p gegen die Richtungen $\alpha_1, \beta_1, \gamma_1; \alpha_2, \beta_2, \gamma_2$ abhängt.

Von den Relationen (14) ausgehend, erhalten wir folgende Ausdrücke für R, S, T allein durch die Funktion $H = H(\vartheta, \psi)$:

$$(17) \quad \begin{cases} R = \frac{\partial^2 H}{\partial \vartheta^2} + H, \\ S = \frac{1}{\sin \vartheta} \frac{\partial^2 H}{\partial \vartheta \partial \psi} - \frac{\cos \vartheta}{\sin^2 \vartheta} \frac{\partial H}{\partial \psi}, \\ T = \frac{1}{\sin^2 \vartheta} \frac{\partial^2 H}{\partial \psi^2} + \frac{\cos \vartheta}{\sin \vartheta} \frac{\partial H}{\partial \vartheta} + H. \end{cases}$$

Dabei wird immer $R > 0$, $RT - S^2 > 0$, und in diesen Ungleichungen sind bereits die allgemeinen Bedingungen (1)–(4) für die Stützebenenfunktion H völlig eingeschlossen.

Setzen wir die Determinante $\left| x, \frac{\partial x}{\partial \vartheta}, \frac{\partial x}{\partial \psi} \right|$ zu ihrer linken Seite mit der Determinante 1 der Substitution (11) zusammen, so finden wir sie $= H(RT - S^2) \sin \vartheta$, so daß aus (15):

$$(18) \quad df = (RT - S^2) d\omega$$

und aus (16):

$$(19) \quad V = \frac{1}{3} \int H(RT - S^2) d\omega$$

hervorgeht.

Das Volumen V erscheint hiernach als ein gewisser homogener Ausdruck dritten Grades in den Werten $H(\vartheta, \psi)$.

15. Ist ein konvexer Körper $\mathfrak{K}$ die Grenze einer unendlichen Reihe vollkommener Ovaloide $\mathfrak{N}', \mathfrak{N}'', \dots$, so konvergieren die Volumina dieser Ovaloide nach dem Volumen von $\mathfrak{K}$. Man erschließt diese Tatsache ganz allein mit Hilfe des Umstandes, daß, wenn ein Ovaloid ein anderes in sich enthält, das erstere stets ein größeres Volumen besitzt. Ferner konvergieren die Koordinaten der Schwerpunkte von $\mathfrak{N}', \mathfrak{N}'', \dots$ nach den Koordinaten des Schwerpunktes von $\mathfrak{K}$.

§ 4. Scharen konvexer Körper. Gemischtes Volumen dreier Körper.

16. Sind $H_1, H_2, \dots, H_m$ die Stützebenenfunktionen von m konvexen Körpern $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$, so genügt die Funktion

$$H(u, v, w) = t_1 H_1(u, v, w) + t_2 H_2(u, v, w) + \dots + t_m H_m(u, v, w),$$

wenn die Parameter $t_1, t_2, \dots, t_m$ sämtlich ≥ 0 , aber nicht durchweg $= 0$ sind, stets ebenfalls allen Bedingungen (1)–(4) in § 1 und bildet daher wiederum die Stützebenenfunktion eines konvexen Körpers. Diesen Körper bezeichnen wir durch

$$(20) \quad \mathfrak{R} = t_1 \mathfrak{R}_1 + t_2 \mathfrak{R}_2 + \dots + t_m \mathfrak{R}_m$$

und die Gesamtheit aller in solcher Weise aus gegebenen Grundkörpern $\mathfrak{R}_i$ herzuleitenden Körper $\mathfrak{R}$ nennen wir eine *Schar konvexer Körper*.

17. Sind $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ *vollkommene Ovaloide*, so gilt das gleiche von jedem Körper $\mathfrak{R}$ der Schar (20). Bezeichnen wir die Punkte auf den Begrenzungen von $\mathfrak{R}, \mathfrak{R}_i$, in welchen eine bestimmte (und die nämliche) äußere Normale α, β, γ vorhanden ist, mit $x, y, z; x_i, y_i, z_i$, so finden wir auf Grund der Gleichungen (14) die Beziehungen gültig:

$$(21) \quad x = t_1 x_1 + t_2 x_2 + \dots + t_n x_n, \dots$$

Stellen wir nun das Volumen V von $\mathfrak{R}$ gemäß der Formel (16) dar, so resultiert mit Rücksicht auf diese Gleichungen (21) für V ein homogener Ausdruck dritten Grades in den Parametern t_i :

$$(22) \quad V = \sum V_{jkl} t_j t_k t_l \quad (j, k, l = 1, 2, \dots, m);$$

die Koeffizienten V_{jkl} mit nicht lauter gleichen Indizes denken wir uns dabei so eingeführt, daß sie bei Permutationen der Indizes sich nicht ändern. Ein jeder Koeffizient V_{jkl} ist allein von den drei zugehörigen Körpern $\mathfrak{R}_j, \mathfrak{R}_k, \mathfrak{R}_l$ abhängig; wir bezeichnen ihn auch mit $V(\mathfrak{R}_j, \mathfrak{R}_k, \mathfrak{R}_l)$ und nennen ihn das *gemischte Volumen* der Körper $\mathfrak{R}_j, \mathfrak{R}_k, \mathfrak{R}_l$.

18. Betrachten wir nun das gemischte Volumen $V_{123} = V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ dreier vollkommener Ovaloide $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$. Schreiben wir

$$J_{123} = \frac{1}{3} \iint \left| x_1, \frac{\partial x_2}{\partial \vartheta}, \frac{\partial x_3}{\partial \psi} \right| d\vartheta d\psi$$

und definieren die analogen Ausdrücke für die Permutationen der Indizes 1, 2, 3, so ist

$$6 V_{123} = (J_{123} + J_{132}) + (J_{231} + J_{213}) + (J_{312} + J_{321}).$$

Nun gilt

$$\frac{1}{3} \iint \frac{\partial}{\partial \vartheta} \left| x_1, x_2, \frac{\partial x_3}{\partial \psi} \right| d\vartheta d\psi = J_{123} - J_{213} + \frac{1}{3} \iint \left| x_1, x_2, \frac{\partial^2 x_3}{\partial \vartheta \partial \psi} \right| d\vartheta d\psi,$$

$$\frac{1}{3} \iint \frac{\partial}{\partial \psi} \left| x_1, x_2, \frac{\partial x_3}{\partial \vartheta} \right| d\vartheta d\psi = J_{231} - J_{132} + \frac{1}{3} \iint \left| x_1, x_2, \frac{\partial^2 x_3}{\partial \vartheta \partial \psi} \right| d\vartheta d\psi.$$

Die linken Seiten in diesen zwei Gleichungen sind gleich Null, weil die Integranden Differentialquotienten nach ϑ , bzw. ψ sind und die Integrationen sich über die ganze Kugelfläche $\mathfrak{E}$ erstrecken. Durch Subtraktion der damit hervorgehenden Relationen folgt nun

$$(23) \quad J_{123} + J_{132} = J_{231} + J_{213},$$

und durch zyklische Vertauschung von 1, 2, 3 hier finden wir weiter diese Ausdrücke $= J_{312} + J_{321}$, so daß sich

$$V_{123} = \frac{1}{2} (J_{123} + J_{132})$$

herausstellt.

Multiplizieren wir in J_{123} die Determinante $\left| x_1, \frac{\partial x_2}{\partial \vartheta}, \frac{\partial x_3}{\partial \psi} \right|$ zur linken Seite mit der Determinante 1 der Substitution (11), und bezeichnen wir die den Formeln (17) gemäß herzustellenden Ausdrücke R, S, T für $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ mit den entsprechenden Indizes, so entsteht

$$J_{123} = \frac{1}{3} \int H_1 (R_2 T_3 - S_2 S_3) d\omega;$$

analog drücken wir J_{132} aus, und wir erhalten endlich

$$(24) \quad V_{123} = \frac{1}{6} \int H_1 (R_2 T_3 - 2 S_2 S_3 + T_2 R_3) d\omega.$$

Nach der Gleichung (23) besteht für diesen Ausdruck die wichtige Eigenschaft, daß er bei beliebigen Permutationen von 1, 2, 3 seinen Wert nicht ändert.

19. Wir knüpfen an diese Formel einige einfache Bemerkungen an. Die Größen R_3, S_3, T_3 bleiben bei einer Translation von $\mathfrak{R}_3$ ungeändert, mithin bleibt dabei auch V_{123} ungeändert, und da $V_{123} = V_{231} = V_{132}$ ist, so folgt allgemein:

Der Wert $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ bleibt bei beliebigen Translationen der einzelnen Körper $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ ungeändert.

Der Ausdruck $R_2 T_3 - 2 S_2 S_3 + T_2 R_3$ ist als die Simultaninvariante zweier positiver binärer quadratischer Formen stets > 0 . Hat $\mathfrak{R}_1$ den Nullpunkt im Inneren liegen (was sich stets durch eine Translation von $\mathfrak{R}_1$ hervorrufen läßt), so ist durchweg $H_1(\vartheta, \psi) > 0$ und also dann auch $V_{123} > 0$. Mit Rücksicht auf den eben bewiesenen Satz gilt daher allgemein:

Die Größe $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ ist stets > 0 .

Ist der Körper $\mathfrak{R}_1$ in einem anderen vollkommenen Ovaloide $\mathfrak{R}_1^*$ mit der Stützebenenfunktion H_1^* enthalten, so gilt stets $H_1(\vartheta, \psi) \leq H_1^*(\vartheta, \psi)$ und entnehmen wir aus (24):

$$(25) \quad V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3) \leq V(\mathfrak{R}_1^*, \mathfrak{R}_2, \mathfrak{R}_3).$$

Endlich bemerken wir die Regel: Ist t ein positiver Faktor, so gilt

$$(26) \quad V(t\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3) = t V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3).$$

20. Sind die Körper $\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3$ mit einem und demselben Körper $\mathfrak{R}$ identisch, so wird $V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_3)$ das *Volumen* von $\mathfrak{R}$.

Die Stützebenenfunktion der Kugel $\mathfrak{G}(x^2 + y^2 + z^2 \leq 1)$ ist für die Systeme α, β, γ auf $\mathfrak{G}$ konstant $= 1$. Beziehen sich H, R, S, T auf ein vollkommenes Ovaloid $\mathfrak{R}$ und bezeichnen wir das Oberflächenelement dieses Körpers mit df , seine ganze *Oberfläche* mit O , so folgt daher aus (24):

$$(27) \quad 3 V(\mathfrak{G}, \mathfrak{R}, \mathfrak{R}) = \int (RT - S^2) d\omega = \int df = O,$$

und durch (23) gelangen wir noch zu dem Ausdrucke

$$(28) \quad O = 3 V(\mathfrak{R}, \mathfrak{G}, \mathfrak{R}) = \frac{1}{2} \int H(R + T) d\omega.$$

Andererseits wird

$$(29) \quad 3 V(\mathfrak{G}, \mathfrak{G}, \mathfrak{R}) = \frac{1}{2} \int (R + T) d\omega = \int \frac{\frac{1}{2}(R + T)}{RT - S^2} df;$$

die Größe $\frac{\frac{1}{2}(R + T)}{RT - S^2}$ hier bedeutet die mittlere Krümmung am Flächenelement df und können wir danach $3 V(\mathfrak{G}, \mathfrak{G}, \mathfrak{R})$ als *das Integral der mittleren Krümmung* von $\mathfrak{R}$ bezeichnen. Wir haben dann noch die Beziehung

$$(30) \quad 3 V(\mathfrak{G}, \mathfrak{G}, \mathfrak{R}) = 3 V(\mathfrak{R}, \mathfrak{G}, \mathfrak{G}) = \int H d\omega.$$

21. Sind $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ beliebige konvexe Körper, so können wir nach § 2 jeden dieser Körper $\mathfrak{R}_i$ als Grenze einer Reihe vollkommener Ovaloide $\mathfrak{D}_i$ darstellen. Auf Grund der in 19. abgeleiteten Regeln läßt sich dann zeigen, daß dabei ein jeder Ausdruck $V(\mathfrak{D}_j, \mathfrak{D}_k, \mathfrak{D}_l)$ immer nach einer bestimmten, von der Wahl der annähernden Ovaloide unabhängigen Grenze konvergiert, die wir mit $V(\mathfrak{R}_j, \mathfrak{R}_k, \mathfrak{R}_l)$ bezeichnen und das *gemischte Volumen* von $\mathfrak{R}_j, \mathfrak{R}_k, \mathfrak{R}_l$ nennen. Es übertragen sich dann alle Regeln aus 19. und die Entwicklung (22) sofort auf beliebige konvexe Körper.

Für die gemischten Volumina bestehen einige fundamentale Ungleichungen, die in dem Satze gipfeln, daß irgend drei Körper vom Volumen 1 stets ein gemischtes Volumen ≥ 1 ergeben. Als Vorbereitung zum Nachweis dieser Ungleichungen behandeln wir zunächst die entsprechenden Fragen für die Ebene.

§ 5. Ovale. Gemischter Flächeninhalt zweier Ovale.

22. Wir betrachten jetzt Figuren in einer Ebene $z = \text{const.}$ Es sei $H(u, v)$ eine reelle Funktion zweier reeller Argumente mit den Eigenschaften:

$$H(0, 0) = 0, \quad H(tu, tv) = tH(u, v), \quad \text{wenn } t > 0 \text{ ist,}$$

$$H(u_1 + u_2, v_1 + v_2) \leq H(u_1, v_1) + H(u_2, v_2),$$

$$H(u, v) + H(-u, -v) > 0, \quad \text{wenn } u, v \neq 0, 0 \text{ ist.}$$

Den durch die sämtlichen Ungleichungen

$$ux + vy \leq H(u, v)$$

für alle möglichen Systeme u, v definierten Bereich $\mathfrak{F}$ von Punkten x, y in der Ebene $z = \text{const.}$ bezeichnen wir als ein *Oval* in dieser Ebene, und wir nennen $H(u, v)$ die *Stützgeradenfunktion* dieses Ovals. Jedem solchen Oval $\mathfrak{F}$ kommt ein bestimmter Flächeninhalt, ferner ein bestimmter Schwerpunkt zu; der Schwerpunkt wird stets ein innerer Punkt des Ovals.

Ist s ein positiver Wert > 0 , so stellt $sH(u, v)$ wieder die Stützgeradenfunktion eines Ovals vor, das wir dann $s\mathfrak{F}$ nennen.

23. Wir bezeichnen ein Oval $\mathfrak{F}$ als ein *vollkommenes Oval*, wenn die Begrenzung von $\mathfrak{F}$ durch eine *analytische* Gleichung in x, y gegeben ist und in jedem Punkte eine *bestimmte* und immer nur eine *Berührung erster Ordnung* eingehende *Tangente* hat.

Ist $\mathfrak{F}$ ein beliebiges Oval mit dem Nullpunkt als Schwerpunkt, und ε eine beliebige positive Größe, so kann man stets ein vollkommenes Oval $\mathfrak{F}^*$, ebenfalls mit dem Nullpunkte als Schwerpunkt, finden, welches $\mathfrak{F}$ enthält und selbst ganz in $(1 + \varepsilon)\mathfrak{F}$ enthalten ist. Jedes Oval kann als *Grenze* vollkommener Ovale mit dem nämlichen Schwerpunkt dargestellt werden.

24. Es sei $\mathfrak{F}$ ein vollkommenes Oval, $H(u, v)$ seine Stützgeradenfunktion. Wir bezeichnen mit α, β die Koordinaten eines Punktes der Kreislinie $x^2 + y^2 = 1$ bzw. die Richtung von $x = 0, y = 0$ nach diesem Punkte, führen $\alpha = \cos \vartheta, \beta = \sin \vartheta, 0 \leq \vartheta \leq 2\pi$ ein und schreiben $H(\alpha, \beta) = H(\vartheta) = H$. Der Krümmungsradius der Begrenzung von $\mathfrak{F}$ in einem Punkte p , wo die äußere Normale die Richtung (α, β) hat, wird $R = H + \frac{\partial^2 H}{\partial \vartheta^2}$ und der Flächeninhalt von $\mathfrak{F}$ bekommt den Ausdruck:

$$(31) \quad F = \frac{1}{2} \int_0^{2\pi} H \left(H + \frac{\partial^2 H}{\partial \vartheta^2} \right) d\vartheta.$$

Wir haben nun in bezug auf den Flächeninhalt $\mathfrak{F}$ noch einige Bemerkungen zu entwickeln, die bald Anwendung finden werden.

Es sei a der kleinste, A der größte Wert von x in $\mathfrak{F}$ und bezeichnen wir für ein $x \geq a$ und $\leq A$ mit $y(x)$ die Länge der zur y -Achse parallelen Sehne von $\mathfrak{F}$, auf welcher der betreffende Abszissenwert x konstant ist. Die Funktion $y(x)$ ist im Inneren des Intervalls $a \leq x \leq A$ regulär und an den Grenzen nähert sie sich stetig dem Werte Null. *Ferner ist darin $\frac{dy}{dx}$ eine mit wachsendem x stets abnehmende Funktion.* Infolgedessen ist weiter insbesondere

$$\frac{\int_a^x \frac{dy}{dx} dx}{\int_a^x dx} = \frac{y}{x-a}$$

eine stets abnehmende und analog $\frac{y}{A-x}$ eine stets wachsende Funktion von x .

Der Flächeninhalt F von $\mathfrak{F}$ besitzt nun auch den Ausdruck

$$(32) \quad F = \int_a^A y(x) dx.$$

Setzen wir, wenn $a \leq x \leq A$ ist,

$$\int_a^x y(x) dx = \tau F,$$

so ist $\tau = \tau(x)$ eine solche Funktion der oberen Grenze x dieses Integrals, welche *kontinuierlich* von 0 bis 1 *zunimmt*, während x von a bis A läuft, und können wir daher *umgekehrt* die obere Grenze dieses Integrals als eine bestimmte Funktion $x(\tau)$ des Wertes τ im Intervalle $0 \leq \tau \leq 1$ einführen. Dabei gilt

$$y \frac{dx}{d\tau} = F, \quad \frac{d(y^2)}{d\tau} = 2F \frac{dy}{dx}.$$

Nach der zweiten Gleichung hier wird $\frac{d(y^2)}{d\tau}$ eine mit wachsendem τ stets abnehmende Funktion von τ ; infolgedessen ist weiter $\frac{y^2}{\tau}$ eine stets abnehmende, $\frac{y^2}{1-\tau}$ eine stets zunehmende Funktion von τ .

Die Länge der Sehne bc von $\mathfrak{F}$ auf der Geraden $x = \frac{a+A}{2}$, also $y\left(\frac{a+A}{2}\right)$, sei $= h$. Die Tangenten an $\mathfrak{F}$ in den Endpunkten dieser Sehne bilden mit den Geraden $x = a$ und $x = A$ ein Trapez, welches $\mathfrak{F}$ ganz in sich enthält; also folgt

$$(33) \quad (A-a)h \geq F.$$

Andererseits enthält $\mathfrak{F}$ das Dreieck abc mit jener Sehne bc als Basis und der Spitze in dem Punkte a von $\mathfrak{F}$, für den $x = a$ ist; daher gilt

$$\frac{1}{4}(A-a)h < \tau \left(\frac{a+A}{2} \right) F,$$

woraus mit Rücksicht auf (33) nun $\tau \left(\frac{a+A}{2} \right) > \frac{1}{4}$ hervorgeht; analog finden wir $\tau \left(\frac{a+A}{2} \right) < \frac{3}{4}$.

Für einen Wert $\tau > \frac{3}{4}$ fällt danach stets $x(\tau) > \frac{a+A}{2}$ aus; da $\frac{y}{A-x}$ und $\frac{y^2}{1-\tau}$ mit wachsendem x bzw. τ zunehmen, haben wir alsdann

$$\frac{y}{A-x} > \frac{h}{\frac{1}{2}(A-a)}, \quad (1-\tau)F = \int_x^A y dx > \frac{(A-x)^2 h}{A-a},$$

woraus mit Rücksicht auf (33) sich

$$(34) \quad A - x(\tau) < (A-a)\sqrt{1-\tau}$$

ergibt, und erhalten wir andererseits

$$(35) \quad \frac{y^2}{1-\tau} > \frac{h^2}{1-\tau \left(\frac{a+A}{2} \right)} > \frac{4}{3} h^2, \quad \frac{y^2}{\tau} > \frac{4}{3} \frac{(1-\tau)}{\tau} \frac{F^2}{(A-a)^2}.$$

Nehmen wir an, der Schwerpunkt von $\mathfrak{F}$ liege im Nullpunkte, so führt die Betrachtung der x -Koordinate des Schwerpunktes zur Gleichung

$$0 = \frac{1}{F} \int_a^A xy dx = \int_0^1 x d\tau,$$

woraus durch partielle Integration

$$(36) \quad A = \int_a^A \tau dx, \quad A = F \int_0^1 \frac{\tau d\tau}{y}$$

folgt. Zerlegen wir $\mathfrak{F}$ in die Dreiecke mit a als Spitze und den einzelnen Bogenelementen des Umfangs von $\mathfrak{F}$ als Grundlinien, so ist für den Schwerpunkt eines jeden dieser Dreiecke offenbar $A - x > \frac{1}{3}(A-a)$ und muß die gleiche Relation daher auch für den Schwerpunkt von $\mathfrak{F}$ selbst gelten, d. h. wir haben

$$(37) \quad A > \frac{A-a}{3}.$$

25. Es seien $\mathfrak{F}_1$ und $\mathfrak{F}_2$ zwei beliebige Ovale und $H_1(u, v)$, $H_2(u, v)$ ihre Stützgeradenfunktionen; alsdann bildet $(1-t)H_1(u, v) + tH_2(u, v)$, wenn $t > 0$ und < 1 ist, immer ebenfalls die Stützgeradenfunktion eines Ovals, das mit $\mathfrak{F} = (1-t)\mathfrak{F}_1 + t\mathfrak{F}_2$ bezeichnet werde. Dieser Herstellung

von $\mathfrak{F}$ steht folgende Erzeugungsweise desselben Ovals dual gegenüber. Man verbinde jeden Punkt f_1 von $\mathfrak{F}_1$ mit jedem Punkte f_2 von $\mathfrak{F}_2$ und teile immer die Verbindungsstrecke $f_1 f_2$ in einem Punkte f so, daß

$$f = (1-t)f_1 + tf_2$$

ist (d. h. daß die Längen der Strecken $f_1 f$ und ff_2 sich wie $t:1-t$ verhalten). Die Menge aller verschiedenen Punkte f , die auf diese Weise gefunden werden, bildet dann genau den Bereich des Ovals $\mathfrak{F}$. Bei dieser Konstruktion von $\mathfrak{F}$ denken wir uns zweckmäßig $\mathfrak{F}_1$ und $\mathfrak{F}_2$ in zwei verschiedenen Ebenen $z = \text{const.}$, etwa $\mathfrak{F}_1$ in $z = 0$ und $\mathfrak{F}_2$ in $z = 1$ gelegen; alsdann stellt $\mathfrak{F} = (1-t)\mathfrak{F}_1 + t\mathfrak{F}_2$ den Schnitt des kleinsten, die beiden Ovale $\mathfrak{F}_1$ und $\mathfrak{F}_2$ ganz enthaltenden konvexen Körpers mit der Ebene $z = t$ vor.

Der Flächeninhalt des Ovals $\mathfrak{F}$ besitzt einen Ausdruck

$$(38) \quad F = (1-t)^2 F_{11} + 2(1-t)t F_{12} + t^2 F_{22},$$

wobei F_{11} den Flächeninhalt von $\mathfrak{F}_1$, F_{22} den von $\mathfrak{F}_2$ und F_{12} eine weitere Konstante bedeutet, die wir den *gemischten Flächeninhalt* von $\mathfrak{F}_1$ und $\mathfrak{F}_2$ nennen. F_{12} ändert sich nicht bei beliebigen Translationen von $\mathfrak{F}_1$ oder von $\mathfrak{F}_2$.

Sind $\mathfrak{F}_1$ und $\mathfrak{F}_2$ vollkommene Ovale, so finden wir, von der Formel (31) ausgehend,

$$F_{12} = \frac{1}{2} \int_0^{2\pi} H_1 \left(H_2 + \frac{\partial^2 H_2}{\partial \vartheta^2} \right) d\vartheta = \frac{1}{2} \int_0^{2\pi} H_2 \left(H_1 + \frac{\partial^2 H_1}{\partial \vartheta^2} \right) d\vartheta,$$

worin H_1 für $H_1(\cos \vartheta, \sin \vartheta)$ und H_2 für $H_2(\cos \vartheta, \sin \vartheta)$ steht.

Stellt $\mathfrak{F}_2$ die Kreisfläche $x^2 + y^2 \leq 1$ vor, so ist H_2 hier konstant $= 1$, $F_{22} = \pi$ und wird $2F_{12}$ die *Länge des Umfangs* von $\mathfrak{F}_1$.

26. Wir wollen nun den Satz beweisen, daß stets die Ungleichung gilt:

$$(39) \quad F_{12} \geq \sqrt{F_{11} F_{22}}.$$

Diese Ungleichung ist gleichbedeutend mit folgender Beziehung für den in (38) entwickelten Ausdruck F :

$$(40) \quad \sqrt{F} \geq (1-t)\sqrt{F_{11}} + t\sqrt{F_{22}}.$$

Von besonderem Werte ist es, auch die Grenzfälle der Ungleichung (39) mit aufzuklären. Wir werden den Zusatz gewinnen, daß in dieser Ungleichung dann und nur dann das Gleichheitszeichen eintritt, wenn $\mathfrak{F}_1$ und $\mathfrak{F}_2$ homothetisch sind (d. h. auseinander durch Translation und Dilation hervorgehen, miteinander ähnlich und ähnlich gelegen sind).

Wir denken uns der Einfachheit wegen den Schwerpunkt von $\mathfrak{F}_1$ wie den von $\mathfrak{F}_2$ im Nullpunkte gelegen (was stets durch Translation dieser Ovale zu erreichen ist). Alsdann sind $H_1(\alpha, \beta)$ und $H_2(\alpha, \beta)$

immer > 0 . Auf der Kreislinie $\alpha^2 + \beta^2 = 1$ hat die daselbst stetige Funktion

$$(41) \quad \frac{\sqrt{F_{11}} H_2(\alpha, \beta)}{\sqrt{F_{22}} H_1(\alpha, \beta)}$$

ein bestimmtes *Maximum*, das wir D nennen. Dieser Wert hat die Bedeutung, daß das Oval $\frac{1}{\sqrt{F_{22}}} \mathfrak{F}_2$ vom Flächeninhalt 1 im Oval $\frac{s}{\sqrt{F_{11}}} \mathfrak{F}_1$ vom Flächeninhalt s^2 enthalten ist, wenn $s \geq D$ ist, aber nicht ganz darin enthalten, wenn $s < D$ ist. Sind $\mathfrak{F}_1$ und $\mathfrak{F}_2$ homothetisch, so decken sich die Ovale $\frac{1}{\sqrt{F_{22}}} \mathfrak{F}_2$ und $\frac{1}{\sqrt{F_{11}}} \mathfrak{F}_1$ und ist daher auch ihr gemischter Flächeninhalt $= 1$, also $F_{12} = \sqrt{F_{11} F_{22}}$ und wird andererseits $D = 1$.

Jetzt verfolgen wir die Annahme, daß $\mathfrak{F}_1$ und $\mathfrak{F}_2$ *nicht homothetisch* sind. Alsdann muß $D > 1$ ausfallen. Wir werden nun eine Ungleichung aufstellen

$$\frac{F_{12}}{\sqrt{F_{11} F_{22}}} - 1 \geq \Pi(D),$$

worin $\Pi(D)$ eine stetige Funktion von D sein wird, die für ein $D > 1$ stets > 0 ist. Diese Beziehung setzt dann in Evidenz, daß stets $F_{12} > \sqrt{F_{11} F_{22}}$ wird, wenn $\mathfrak{F}_1$ und $\mathfrak{F}_2$ nicht homothetisch sind.

Eine Beziehung von diesem Charakter mit einer stetigen Funktion $\Pi(D)$ braucht offenbar nur für vollkommene Ovale erwiesen zu werden. Durch unseren Hilfssatz in 23. über die Annäherung eines beliebigen Ovals durch vollkommene Ovale gilt sie dann sofort für alle Ovale.

27. Es seien nunmehr $\mathfrak{F}_1$ und $\mathfrak{F}_2$ zwei vollkommene Ovale und sie seien nicht homothetisch. Wir drehen die Koordinatenachsen um den Nullpunkt derart, daß die positive x -Achse in eine solche Richtung fällt, wofür die Funktion (41) ihren größten Wert D erhält, d. h. also, wir nehmen

$$(42) \quad \frac{\sqrt{F_{11}} H_2(1, 0)}{\sqrt{F_{22}} H_1(1, 0)} = D$$

an.

Es sei jetzt t irgendein bestimmter Wert > 0 und < 1 ; wir gebrauchen für das Oval $\mathfrak{F} = (1-t)\mathfrak{F}_1 + t\mathfrak{F}_2$ die Zeichen $a, A, y(x)$ in derselben Weise, wie sie für ein Oval $\mathfrak{F}$ in 24. erklärt sind; die entsprechenden Größen für $\mathfrak{F}_1$ und $\mathfrak{F}_2$ bezeichnen wir mit $a_1, A_1, y_1(x)$, $a_2, A_2, y_2(x)$; insbesondere ist $A_1 = H_1(1, 0)$, $A_2 = H_2(1, 0)$. Alsdann bestehen für a und A , den kleinsten bzw. größten Wert von x in $\mathfrak{F}$, die Ausdrücke

$$(43) \quad a = (1-t)a_1 + ta_2, \quad A = (1-t)A_1 + tA_2.$$

Die Flächeninhalte von $\mathfrak{F}_1$ und $\mathfrak{F}_2$ stellen wir gemäß der Formel (32) dar:

$$(44) \quad F_{11} = \int_{a_1}^{A_1} y_1(x) dx, \quad F_{22} = \int_{a_2}^{A_2} y_2(x) dx.$$

Wir setzen nun, wenn $a_1 \leq x_1 \leq A_1$, $a_2 \leq x_2 \leq A_2$ ist,

$$(45) \quad \int_{a_1}^{x_1} y_1(x) dx = \tau F_{11}, \quad \int_{a_2}^{x_2} y_2(x) dx = \tau F_{22}$$

und führen die oberen Grenzen dieser Integrale als Funktionen $x_1(\tau)$, $x_2(\tau)$ des Wertes τ ein, der sich im Intervalle $0 \leq \tau \leq 1$ bewegt. Indem wir unter τ in diesen beiden Funktionen ein und dasselbe Argument verstehen, wird ein gewisser Zusammenhang zwischen den zur y -Achse parallelen Sehnen von $\mathfrak{F}_1$ und von $\mathfrak{F}_2$ hergestellt.*)

Wir ziehen jetzt die in 25. angegebene Methode zur Erzeugung des Ovals $\mathfrak{F}$ aus den Punkten von $\mathfrak{F}_1$ und $\mathfrak{F}_2$ heran und bringen sie zunächst auf die Punkte der Sehne $p_1 q_1$ von $\mathfrak{F}_1$ auf der Geraden $x = x_1(\tau)$ und der Sehne $p_2 q_2$ von $\mathfrak{F}_2$ auf der Geraden $x = x_2(\tau)$ in Anwendung; dadurch erkennen wir, daß zu $\mathfrak{F}$ auf der Geraden

$$(46) \quad x = (1 - t)x_1(\tau) + tx_2(\tau)$$

jedenfalls alle diejenigen Punkte gehören müssen, welche diese Gerade aus dem Trapeze $p_1 q_1 q_2 p_2$ herausschneidet. Die Längen der Sehnen $p_1 q_1$ und $p_2 q_2$ sind $y_1(x_1(\tau))$ bzw. $y_2(x_2(\tau))$; danach muß für die Länge $y(x)$ der Sehne $p q$ von $\mathfrak{F}$ auf der durch (46) angewiesenen Geraden jedenfalls die Ungleichung

$$(47) \quad y(x) \geq (1 - t)y_1(x_1(\tau)) + ty_2(x_2(\tau))$$

gelten.

Für den Flächeninhalt F des Ovals $\mathfrak{F}$ haben wir

$$F = \int_a^A y(x) dx.$$

Nun wächst die durch (46) gegebene Funktion x kontinuierlich und läuft nach (43) in den Grenzen a bis A , wenn τ von 0 bis 1 zunimmt, und können wir sie daher als Variable in dieses Integral für F einführen. Schreiben wir zur Abkürzung x_1 , y_1 und x_2 , y_2 für $x_1(\tau)$, $y_1(x_1(\tau))$ und $x_2(\tau)$, $y_2(x_2(\tau))$ und machen von (47) Gebrauch, so ergibt sich

*) Die Idee dieser Zuordnung der parallelen Sehnen in zwei Ovalen nach dem Verhältnis der je zwei Flächenstücke, in welche sie die Ovale zerschneiden, rührt von Herrn H. Brunn her (vgl. *Ovale und Eiflächen*, Inauguraldissertation, München, 1887, S. 23).

$$F \geq \int_0^1 ((1-t)y_1 + ty_2) \left((1-t) \frac{dx_1}{d\tau} + t \frac{dx_2}{d\tau} \right) d\tau.$$

Jetzt ziehen wir den Ausdruck (38) für F heran und benutzen die Gleichungen (44); dadurch erhalten wir

$$2F_{12} \geq \int_0^1 \left(y_1 \frac{dx_2}{d\tau} + y_2 \frac{dx_1}{d\tau} \right) d\tau.$$

Hier führen wir gemäß (45):

$$\frac{dx_1}{d\tau} = \frac{F_{11}}{y_1}, \quad \frac{dx_2}{d\tau} = \frac{F_{22}}{y_2}$$

ein, und wir erzielen endlich

$$(48) \quad 2(F_{12} - \sqrt{F_{11}F_{22}}) \geq \int_0^1 \frac{(y_1 \sqrt{F_{22}} - y_2 \sqrt{F_{11}})^2}{y_1 y_2} d\tau.$$

Andererseits haben wir der Formel (36) zufolge:

$$A_1 = F_{11} \int_0^1 \frac{\tau d\tau}{y_1}, \quad A_2 = F_{22} \int_0^1 \frac{\tau d\tau}{y_2},$$

mithin

$$(49) \quad \int_0^1 \frac{y_1 \sqrt{F_{22}} - y_2 \sqrt{F_{11}}}{y_1 y_2} \tau d\tau = \frac{A_2}{\sqrt{F_{22}}} - \frac{A_1}{\sqrt{F_{11}}} = (D-1) \frac{A_1}{\sqrt{F_{11}}} > 0.$$

Mit Hilfe der letzteren Beziehung suchen wir eine positive untere Grenze für das Integral auf der rechten Seite in (48) herzuleiten; dabei entsteht eine Schwierigkeit durch den Umstand, daß $\frac{\tau}{\sqrt{y_1 y_2}}$ bei Annäherung an $\tau = 1$ über jede Grenze hinauswächst. Es sei nun δ eine positive Größe $< \frac{1}{4}$, so gilt

$$F_{22} \int_{1-\delta}^1 \frac{\tau d\tau}{y_2} < F_{22} \int_{1-\delta}^1 \frac{d\tau}{y_2} = A_2 - x_2 (1-\delta);$$

mit Rücksicht auf (34) und (37) finden wir die rechte Seite hier $< 3A_2 \sqrt{\delta}$, und wir erhalten aus (49) die Ungleichung

$$(50) \quad \int_0^{1-\delta} \frac{y_1 \sqrt{F_{22}} - y_2 \sqrt{F_{11}}}{y_1 y_2} \tau d\tau > (D(1-3\sqrt{\delta}) - 1) \frac{A_1}{\sqrt{F_{11}}}.$$

Die Ungleichung (48) bleibt bestehen, wenn wir die Integration rechts nur bis zur oberen Grenze $1-\delta$ erstrecken. Im Intervalle 0 bis $1-\delta$ sind zufolge einer in 24. gemachten Bemerkung $\frac{y_1}{\sqrt{\tau}}$ und $\frac{y_2}{\sqrt{\tau}}$

beständig abnehmende Funktionen und haben wir mit Rücksicht auf (35) und (37) und auf $\tau < 1$ durchweg

$$\frac{y_1 y_2}{\tau^2} > \frac{4}{27} \delta \frac{F_{11} F_{22}}{A_1 A_2};$$

damit erhalten wir aus (48) die Beziehung

$$2(F_{12} - \sqrt{F_{11} F_{22}}) \geq \frac{4}{27} \delta \frac{F_{11} F_{22}}{A_1 A_2} \int_0^{1-\delta} \frac{(y_1 \sqrt{F_{22}} - y_2 \sqrt{F_{11}})^2}{y_1^2 y_2^2} \tau^2 d\tau.$$

Bezeichnen wir den Integranden in (50) zur Abkürzung mit Ψ , so ist

$$\int_0^{1-\delta} (\Psi + s)^2 d\tau$$

für jeden reellen Wert von s stets ≥ 0 und gilt infolgedessen

$$(51) \quad \int_0^{1-\delta} \Psi^2 d\tau \geq \frac{1}{1-\delta} \left(\int_0^{1-\delta} \Psi d\tau \right)^2.$$

Diese Beziehung führt uns nun mit Rücksicht auf (50) und, wenn wir noch hier rechts $1 - \delta$ im Nenner durch 1 ersetzen, zu

$$2(F_{12} - \sqrt{F_{11} F_{22}}) \geq \frac{4}{27} \delta \frac{F_{11} F_{22}}{A_1 A_2} \frac{A_1^2}{F_{11}} (D(1 - 3\sqrt{\delta}) - 1)^2.$$

Das Maximum von $3\sqrt{\delta}(D - 1 - 3D\sqrt{\delta})$ tritt für $3\sqrt{\delta} = \frac{D-1}{2D}$ ein, wobei $\delta < \frac{1}{36}$, also gewiß $< \frac{1}{4}$ ist, und wird $= \frac{(D-1)^2}{4D}$. Mit diesem Werte von δ entsteht, wenn wir noch von (42) Gebrauch machen:

$$(52) \quad \frac{F_{12}}{\sqrt{F_{11} F_{22}}} - 1 \geq \frac{1}{2^3 \cdot 3^5} \frac{(D-1)^4}{D^3}.$$

In der hier geschriebenen Form (mit dem Zeichen $\geq$) überträgt sich diese Relation unmittelbar auf beliebige Ovale; sie liefert in der Tat den Satz, daß stets $F_{12} > \sqrt{F_{11} F_{22}}$ ist, wenn $\mathfrak{F}_1$ und $\mathfrak{F}_2$ nicht homothetisch sind.

28. Nehmen wir für $\mathfrak{F}_2$ eine Kreisfläche vom Radius 1, so bedeutet $2F_{12}$ die Länge des Umfangs von $\mathfrak{F}_1$, während $F_{22} = \pi$ zu setzen ist. Die für diesen Fall aus (39) entstehende Ungleichung

$$\frac{2F_{12}}{2\pi} \geq \sqrt{\frac{F_{11}}{\pi}}$$

besagt, daß jedes Oval, welches von einem Kreise verschieden ist, immer kleineren Inhalt besitzt als ein Kreis von demselben Umfange. Diese Aussage läßt sich sogleich auch auf nicht konvexe Flächen ausdehnen, indem zu jeder abgeschlossenen und zusammenhängenden Figur, welche nicht ein Oval vorstellt und welcher ein bestimmter Flächeninhalt und eine

bestimmte Länge des Umfanges zukommt, immer ein bestimmtes kleinstes, die Figur ganz enthaltendes Oval gehört, und für dieses der Flächeninhalt größer, der Umfang kleiner ausfällt als für die Ausgangsfigur.

§ 6. Eine kubische Ungleichung für die gemischten Volumina.

29. Wir suchen jetzt die entsprechenden Tatsachen für den Raum abzuleiten. Zunächst schicken wir einige Hilfsbemerkungen voraus.

Es sei $\mathfrak{R}$ ein vollkommenes Ovaloid, c der kleinste, C der größte Wert von z in $\mathfrak{R}$; für jeden Wert $z \geq c$ und $\leq C$ bezeichnen wir mit $\mathfrak{F}(z)$ den Schnitt von $\mathfrak{R}$ mit der Ebene, in welcher der betreffende Wert z konstant ist, mit $(f(z))^2$ den Flächeninhalt von $\mathfrak{F}(z)$. Für die Werte im Inneren des Intervalls $c \leq z \leq C$ stellt $\mathfrak{F}(z)$ Ovale vor und ist die Funktion $f(z)$ immer regulär, an den Grenzen nähert sich $f(z)$ stetig dem Werte Null. Ist $c < z' < z < z'' < C$, $z = (1-t)z' + tz''$, so enthält das Schnitt-oval $\mathfrak{F}(z)$, da $\mathfrak{R}$ ein konvexer Körper ist, jedenfalls das ganze durch $(1-t)\mathfrak{F}(z') + t\mathfrak{F}(z'')$ dargestellte Oval in sich und gilt daher zufolge der Ungleichung (40) sicherlich

$$f(z) \geq (1-t)f(z') + tf(z'').$$

Auf Grund dieser Eigenschaften erkennen wir, wenn wir z und f als rechtwinklige Koordinaten in einer Ebene deuten, daß der Bereich

$$(53) \quad c \leq z \leq C, \quad 0 \leq f \leq f(z)$$

ein Oval in dieser Ebene vorstellt, und nimmt infolgedessen $\frac{df(z)}{dz}$ mit wachsendem z beständig ab.

Das Volumen V von $\mathfrak{R}$ drückt sich durch

$$(54) \quad V = \int_c^C (f(z))^2 dz$$

aus. Setzen wir nun, wenn $c \leq z \leq C$ ist,

$$\int_c^z (f(z))^2 dz = \tau V,$$

so wächst $\tau = \tau(z)$ kontinuierlich von 0 bis 1, während die obere Grenze z von c bis C zunimmt; wir können dann umgekehrt die obere Grenze z dieses Integrals als eine bestimmte Funktion $z(\tau)$ des Wertes τ einführen, die ihrerseits kontinuierlich von c bis C zunehmen wird, während τ von 0 bis 1 wächst. Dabei gilt, wenn wir zur Abkürzung $f(z(\tau)) = f$ setzen,

$$f^2 \frac{dz}{d\tau} = V.$$

Wir entnehmen daraus weiter

$$\frac{d(f^3)}{d\tau} = 3 \frac{df}{dz} V,$$

so daß auch $\frac{d(f^3)}{d\tau}$ mit wachsendem τ beständig abnimmt; infolgedessen wird weiter $\frac{f^3}{\tau}$ eine beständig abnehmende und andererseits $\frac{f^3}{1-\tau}$ eine beständig zunehmende Funktion von τ sein.

Betrachten wir wieder das durch (53) bestimmte Oval in seiner zf -Ebene. Es sei h die Länge derjenigen Sehne dieses Ovals, auf der $z = \frac{c+C}{2}$ ist, also $h = f\left(\frac{c+C}{2}\right)$. Die Tangente des Ovals im Endpunkte $f = h$ dieser Sehne bildet mit den Geraden $z = c$, $z = C$ und $f = 0$ ein Trapez, in welchem das Oval ganz enthalten ist; dadurch ergibt sich

$$(55) \quad V < \frac{4}{3} (C - c) h^2.$$

Andererseits ist das Dreieck mit jener Sehne als Basis und der Spitze $z = c$, $f = 0$ ganz im Bereiche $z \leq \frac{c+C}{2}$ des Ovals enthalten, und geht hieraus

$$\frac{1}{6} (C - c) h^2 < \tau \left(\frac{c+C}{2}\right) V$$

hervor; mit Rücksicht auf die Relation (55) erhalten wir sodann

$$\tau \left(\frac{c+C}{2}\right) > \frac{1}{8}.$$

Analog finden wir $1 - \tau \left(\frac{c+C}{2}\right) < \frac{1}{8}$, also $\tau \left(\frac{c+C}{2}\right) > \frac{7}{8}$.

Für einen Wert $\tau > \frac{7}{8}$ wird demnach immer $z(\tau) > \frac{c+C}{2}$ sein; es folgt dann

$$\frac{f(z)}{C-z} > \frac{h}{\frac{1}{2}(C-c)}, \quad (1-\tau) V = \int_z^C (f(z))^2 dz > \frac{4}{3} \frac{(C-z)^3}{(C-c)^3} h^2,$$

woraus mit Hilfe von (55) die Ungleichung

$$(56) \quad C - z(\tau) < (C - c) \sqrt[3]{1 - \tau}$$

hervorgeht, und es ergibt sich ferner

$$(57) \quad \frac{(f(z))^3}{1-\tau} > \frac{h^3}{1-\tau \left(\frac{c+C}{2}\right)} > \frac{8}{7} h^3, \quad \frac{(f(z))^3}{\tau} > \frac{8}{7} \frac{1-\tau}{\tau} \left(\frac{3}{4} \frac{V}{C-c}\right)^{\frac{3}{2}}.$$

Nehmen wir den Schwerpunkt von $\mathfrak{R}$ im Nullpunkt gelegen an, so führt die Betrachtung der z -Koordinate dieses Schwerpunkts zur Gleichung

$$(58) \quad 0 = \frac{1}{V} \int_c^C z f^2 dz = \int_0^1 z d\tau, \quad C = \int_c^C \tau dz = V \int_0^1 \frac{\tau d\tau}{f^2}.$$

Zerlegen wir $\mathfrak{R}$ in lauter Pyramiden mit der Spitze in demjenigen Punkte von $\mathfrak{R}$, für den $z = c$ ist, und mit den einzelnen Oberflächenelementen von $\mathfrak{R}$ als Grundflächen, so ist für den Schwerpunkt einer jeden dieser Pyramiden $C - z > \frac{1}{4}(C - c)$ und gilt daher die gleiche Relation auch für den Schwerpunkt von $\mathfrak{R}$ selbst, d. h. man hat

$$(59) \quad C - c < 4C.$$

30. Es seien jetzt $\mathfrak{R}_1$ und $\mathfrak{R}_2$ zwei beliebige konvexe Körper mit den Stützebenenfunktionen H_1 und H_2 und t ein beliebiger Wert > 0 und < 1 . Verbindet man jeden Punkt $\mathfrak{f}_1$ von $\mathfrak{R}_1$ mit jedem Punkte $\mathfrak{f}_2$ von $\mathfrak{R}_2$ und teilt die Strecke $\mathfrak{f}_1\mathfrak{f}_2$ immer in dem Punkte $\mathfrak{f} = (1 - t)\mathfrak{f}_1 + t\mathfrak{f}_2$, so erfüllt die Gesamtheit aller verschiedenen in dieser Weise entstehenden Punkte $\mathfrak{f}$ genau den konvexen Körper $\mathfrak{R} = (1 - t)\mathfrak{R}_1 + t\mathfrak{R}_2$, der als Stützebenenfunktion $H = (1 - t)H_1 + tH_2$ besitzt (vgl. hierzu die Formel (21) in § 4). Das Volumen dieses Körpers $\mathfrak{R}$ wird durch einen Ausdruck

$$(60) \quad V = (1 - t)^3 V_0 + 3(1 - t)^2 t V_1 + 3(1 - t)t^2 V_2 + t^3 V_3$$

dargestellt, worin V_0, V_1, V_2, V_3 die gemischten Volumina

$$(61) \quad V(\mathfrak{R}_1, \mathfrak{R}_1, \mathfrak{R}_1), \quad V(\mathfrak{R}_1, \mathfrak{R}_1, \mathfrak{R}_2), \quad V(\mathfrak{R}_1, \mathfrak{R}_2, \mathfrak{R}_2), \quad V(\mathfrak{R}_2, \mathfrak{R}_2, \mathfrak{R}_2)$$

bedeuten, insbesondere also V_0 das Volumen von $\mathfrak{R}_1$ und V_3 das von $\mathfrak{R}_2$ vorstellt. Diese vier Konstanten bleiben bei beliebigen Translationen von $\mathfrak{R}_1$ und von $\mathfrak{R}_2$ ungeändert.

Wir wollen nun den Satz beweisen:

Für diese Konstanten im Ausdruck von V gilt stets die Ungleichung

$$(62) \quad V_1^3 \geq V_0^2 V_3$$

und zwar tritt hier das Gleichheitszeichen dann und nur dann ein, wenn $\mathfrak{R}_1$ und $\mathfrak{R}_2$ homothetisch sind (auseinander durch Translation und Dilatation hervorgehen).

Wir denken uns von vornherein mit $\mathfrak{R}_1$ und $\mathfrak{R}_2$ solche Translationen vorgenommen, daß sowohl der Schwerpunkt von $\mathfrak{R}_1$ wie der von $\mathfrak{R}_2$ in den Nullpunkt zu liegen kommt. Alsdann sind die Stützebenenfunktionen dieser Körper, $H_1(u, v, w)$ und $H_2(u, v, w)$, für alle Argumente $u, v, w \neq 0, 0, 0$ stets > 0 . Es sei D das Maximum der Funktion

$$(63) \quad \frac{\sqrt[3]{V_0} H_2(\alpha, \beta, \gamma)}{\sqrt[3]{V_3} H_1(\alpha, \beta, \gamma)}$$

auf der ganzen Kugelfläche $\mathfrak{G}$ ($\alpha^2 + \beta^2 + \gamma^2 = 1$); alsdann ist der Körper $\frac{1}{\sqrt[3]{V_3}} \mathfrak{R}_2$ vom Volumen 1 im Körper $\frac{s}{\sqrt[3]{V_0}} \mathfrak{R}_1$ vom Volumen s^3 enthalten, wenn $s \geq D$ ist, aber nicht ganz darin enthalten, wenn $s < D$ ist. Wenn

$\mathfrak{R}_1$ und $\mathfrak{R}_2$ homothetisch sind, gilt $D = 1$ und entnehmen wir aus der Regel (26) in 19., daß alsdann $\frac{V_1}{V_0} = \frac{V_2}{V_1} = \frac{V_3}{V_2}$ und daher $V_1^3 = V_0^2 V_3$ ist.

Sind $\mathfrak{R}_1$ und $\mathfrak{R}_2$ nicht homothetisch, was wir jetzt annehmen wollen, so fällt $D > 1$ aus, und wir werden zeigen, daß alsdann eine Ungleichung besteht:

$$\sqrt[3]{\frac{V_1}{V_0^2 V_3}} - 1 \geq \Pi(D),$$

worin $\Pi(D)$ eine stetige Funktion von D vorstellt, welche für ein $D > 1$ stets > 0 ausfällt. Eine derartige Relation mit einer *stetigen* Funktion $\Pi(D)$ braucht nur für vollkommene Ovaloide $\mathfrak{R}_1, \mathfrak{R}_2$ bewiesen zu werden; vermöge unseres Hilfssatzes aus § 2 über die Annäherung beliebiger konvexer Körper durch vollkommene Ovaloide gilt sie dann sofort für alle konvexen Körper.

31. Es seien jetzt $\mathfrak{R}_1$ und $\mathfrak{R}_2$ zwei vollkommene Ovaloide und sie seien nicht einander homothetisch. Wir geben den Koordinatenachsen eine solche Lage, daß die positive z -Achse in eine Richtung fällt, wofür die Funktion (63) ihren größten Wert D annimmt. Wir verwenden die Zeichen $c, C, \mathfrak{F}(z), f(z)$ für den Körper $\mathfrak{R} = (1-t)\mathfrak{R}_1 + t\mathfrak{R}_2$, wie sie für den Körper $\mathfrak{R}$ in 29. erklärt sind, und bezeichnen die entsprechenden Größen und Bereiche in bezug auf $\mathfrak{R}_1$ oder $\mathfrak{R}_2$ analog unter Hinzufügung des Index 1 oder 2. Als dann ist $C_1 = H_1(1, 0, 0)$, $C_2 = H_2(1, 0, 0)$ und nach der eben gemachten Voraussetzung wird

$$(64) \quad \sqrt[3]{\frac{V_0 C_2}{V_3 C_1}} = D > 1.$$

Der kleinste und größte Wert von z in $\mathfrak{R}$ bestimmen sich gemäß

$$(65) \quad c = (1-t)c_1 + tc_2, \quad C = (1-t)C_1 + tC_2.$$

Wir setzen nun, wenn $c_1 \leq z_1 \leq C_1$, $c_2 \leq z_2 \leq C_2$ ist,

$$(66) \quad \int_{c_1}^{z_1} (f_1(z))^2 dz = \tau V_0, \quad \int_{c_2}^{z_2} (f_2(z))^2 dz = \tau V_3,$$

und führen umgekehrt die oberen Grenzen dieser zwei Integrale als Funktionen $z_1(\tau), z_2(\tau)$ der Variablen τ ein, die sich im Intervalle $0 \leq \tau \leq 1$ bewegt. Indem wir für τ in beiden Funktionen dasselbe Argument verwenden, erhalten wir eine gewisse Zuordnung der Schnitte $\mathfrak{F}_1(z_1(\tau))$ von $\mathfrak{R}_1$ und $\mathfrak{F}_2(z_2(\tau))$ von $\mathfrak{R}_2$.

Andererseits stellen wir das Volumen von $\mathfrak{R}$ in der Formel

$$(67) \quad V = \int_c^C (f(z))^2 dz$$

dar. Setzen wir

$$(68) \quad z = (1 - t)z_1(\tau) + tz_2(\tau),$$

so nimmt diese Funktion von τ , wenn τ von 0 bis 1 läuft, kontinuierlich wachsend die Werte von c bis C an. Wir wenden nun die in 30. erörterte Konstruktion, welche aus den Punkten von $\mathfrak{R}_1$ und von $\mathfrak{R}_2$ die Punkte von $\mathfrak{R}$ herzuleiten gestattet, speziell auf die Punkte f_1 von $\mathfrak{R}_1$ in seinem Schnitte $\mathfrak{F}_1(z_1(\tau))$ und die Punkte f_2 von $\mathfrak{R}_2$ in seinem Schnitte $\mathfrak{F}_2(z_2(\tau))$ an, wobei τ irgendein Wert > 0 und < 1 sei. Die Menge der dabei hervorgehenden Punkte $\bar{f} = (1 - t)f_1 + tf_2$ bildet alsdann genau den Bereich

$$(69) \quad (1 - t)\mathfrak{F}_1(z_1(\tau)) + t\mathfrak{F}_2(z_2(\tau))$$

in derjenigen Ebene $z = \text{const.}$, die durch den Wert z in (68) bestimmt ist; danach muß der zu diesem Werte z gehörige Schnitt $\mathfrak{F}(z)$ von $\mathfrak{R}$ gewiß den ganzen Bereich (69) in sich enthalten. Ziehen wir die Relation (40) aus 26. heran, so geht daraus

$$(70) \quad f(z) \geq (1 - t)f_1(z_1(\tau)) + tf_2(z_2(\tau))$$

hervor. Aus (67) erhalten wir nunmehr, wenn wir für $z_1(\tau)$, $f_1(z_1(\tau))$, $z_2(\tau)$, $f_2(z_2(\tau))$ zur Abkürzung z_1 , f_1 , z_2 , f_2 schreiben:

$$V \geq \int_0^1 ((1 - t)f_1 + tf_2)^2 \left((1 - t) \frac{dz_1}{d\tau} + t \frac{dz_2}{d\tau} \right) d\tau.$$

Wir verwenden jetzt den Ausdruck (60) für V , machen noch von

$$V_0 = \int_0^1 f_1^2 \frac{dz_1}{d\tau} d\tau$$

Gebrauch und lassen endlich in der entstehenden Ungleichung die Größe t sich der Grenze Null nähern; dadurch kommen wir zur Ungleichung

$$3V_1 \geq \int_0^1 \left(f_1^2 \frac{dz_2}{d\tau} + 2f_1f_2 \frac{dz_1}{d\tau} \right) d\tau.$$

Nach den Gleichungen (66) haben wir

$$f_1^2 \frac{dz_1}{d\tau} = V_0, \quad f_2^2 \frac{dz_2}{d\tau} = V_3,$$

und so ergibt sich weiter

$$3V_1 \geq \int_0^1 \left(V_3 \frac{f_1^2}{f_2^2} + 2V_0 \frac{f_2}{f_1} \right) d\tau.$$

Setzen wir

$$(71) \quad \frac{f_1}{\sqrt[3]{V_0}} = \varphi_1, \quad \frac{f_2}{\sqrt[3]{V_3}} = \varphi_2,$$

so läßt sich diese letzte Ungleichung in

$$(72) \quad 3 \left(\sqrt[3]{\frac{V_1}{V_0 \cdot V_3}} - 1 \right) \geq \int_0^1 \left(\frac{\varphi_1^2}{\varphi_2^2} - 3 + \frac{2\varphi_2}{\varphi_1} \right) d\tau = \int_0^1 \frac{(\varphi_1 - \varphi_2)^2 (\varphi_1 + 2\varphi_2)}{\varphi_1 \varphi_2^2} d\tau$$

umwandeln.

Andererseits haben wir gemäß (58):

$$\frac{C_1}{\sqrt[3]{V_0}} = \int_0^1 \frac{\tau d\tau}{\varphi_1^2}, \quad \frac{C_2}{\sqrt[3]{V_3}} = \int_0^1 \frac{\tau d\tau}{\varphi_2^2},$$

mithin

$$(73) \quad \int_0^1 \frac{(\varphi_1 - \varphi_2)(\varphi_1 + \varphi_2)\tau}{\varphi_1^2 \varphi_2^2} d\tau = \frac{C_2}{\sqrt[3]{V_3}} - \frac{C_1}{\sqrt[3]{V_0}} = (D - 1) \frac{C_1}{\sqrt[3]{V_0}} > 0.$$

Nehmen wir jetzt eine positive GröÙe $\delta < \frac{1}{8}$ an, so wird unter Verwendung von (56) und (59):

$$\int_{1-\delta}^1 \frac{\tau d\tau}{\varphi_2^2} < \int_{1-\delta}^1 \frac{d\tau}{\varphi_2^2} = \frac{C_2 - z_2(1-\delta)}{\sqrt[3]{V_3}} < 4\sqrt[3]{\delta} \frac{C_2}{\sqrt[3]{V_3}},$$

und mit Rücksicht hierauf folgt aus (73):

$$(74) \quad \int_0^{1-\delta} \frac{(\varphi_1 - \varphi_2)(\varphi_1 + \varphi_2)\tau}{\varphi_1^2 \varphi_2^2} d\tau > (D(1 - 4\sqrt[3]{\delta}) - 1) \frac{C_1}{\sqrt[3]{V_0}}.$$

Die Ungleichung (72) bleibt gültig, wenn wir auf der rechten Seite die Integration nur bis zur oberen Grenze $1 - \delta$ erstrecken. Bezeichnen wir den Integranden in (74) mit Ψ , so wird der Integrand in (72):

$$(75) \quad \frac{(\varphi_1 + 2\varphi_2) \varphi_1^2 \varphi_2^2}{(\varphi_1 + \varphi_2)^2 \tau^2} \Psi.$$

Nun ist im Intervalle $0 < \tau \leq 1 - \delta$ zufolge (57) und (59):

$$\frac{\varphi_1^2}{\tau^{\frac{2}{3}}} > \frac{3}{7^{\frac{2}{3}} \cdot 4} \delta^{\frac{2}{3}} \frac{\sqrt[3]{V_0}}{C_1}, \quad \frac{\varphi_2^2}{\tau^{\frac{2}{3}}} > \frac{3}{7^{\frac{2}{3}} \cdot 4} \delta^{\frac{2}{3}} \frac{\sqrt[3]{V_3}}{C_2},$$

wobei von den zwei unteren Grenzen hier wegen (64) die erste die größere ist; andererseits hat man $\frac{\varphi_1(\varphi_1 + 2\varphi_2)}{(\varphi_1 + \varphi_2)^2} \geq \frac{3}{4}$, wenn $\varphi_1 \geq \varphi_2$ ist, und $\frac{\varphi_2(\varphi_1 + 2\varphi_2)}{(\varphi_1 + \varphi_2)^2} \geq \frac{3}{4}$, wenn $\varphi_2 \geq \varphi_1$ ist. Danach fällt der Faktor von Ψ in (75) im Intervalle $0 < \tau \leq 1 - \delta$ stets

$$> \frac{27}{7^{\frac{4}{3}} \cdot 64} \frac{\sqrt[3]{V_0 V_3}}{C_1 C_2} \delta^{\frac{4}{3}}$$

aus. Auf Grund der auch schon früher verwandten Hilfsformel (51) erhalten wir nunmehr:

$$3 \left(\sqrt[3]{\frac{V_1}{V_0^2 V_3}} - 1 \right) > \frac{3^3}{7^{\frac{4}{3}} \cdot 2^{14}} (4 \sqrt[3]{\delta})^4 \frac{(D(1 - 4 \sqrt[3]{\delta}) - 1)^2}{D}.$$

Das Produkt $(4 \sqrt[3]{\delta})^2 (D - 1 - 4 \sqrt[3]{\delta} D)$ nimmt seinen größten Wert für $4 \sqrt[3]{\delta} = \frac{2(D-1)}{3D}$ an, wobei $\delta < \frac{1}{216}$, also sicher $< \frac{1}{8}$ ist, und wird dann $= \frac{4(D-1)^3}{27D^2}$. Mit diesem Werte für δ folgt endlich

$$(76) \quad \sqrt[3]{\frac{V_1}{V_0^2 V_3}} - 1 \geq \frac{1}{2^{10} \cdot 3^4 \cdot 7^{\frac{4}{3}}} \frac{(D-1)^6}{D^5}.$$

In solcher Form überträgt sich diese Ungleichung sofort auf zwei beliebige konvexe Körper $\mathfrak{R}_1$ und $\mathfrak{R}_2$ und setzt in der Tat in Evidenz, daß, wenn $\mathfrak{R}_1$ und $\mathfrak{R}_2$ nicht homothetisch sind, stets

$$V_1 > \sqrt[3]{V_0^2 V_3}$$

ausfällt.

32. Vertauschen wir die Rollen von $\mathfrak{R}_1$ und $\mathfrak{R}_2$, so treten, wie aus (61) zu ersehen ist, an Stelle von V_0, V_1, V_2, V_3 diese Größen in umgekehrter Folge und geht aus $V_1 \geq \sqrt[3]{V_0^2 V_3}$ die Ungleichung $V_2 \geq \sqrt[3]{V_0 V_3^2}$ hervor. Diese zwei Ungleichungen zusammen ergeben für den Ausdruck V in (60) die Abschätzung

$$(77) \quad \sqrt[3]{V} \geq (1-t) \sqrt[3]{V_0} + t \sqrt[3]{V_3}.$$

Bezeichnen wir das Minimum der Funktion (63) auf der Kugelfläche $\mathfrak{E}$ mit d , so ist der Körper $\sqrt[3]{\frac{d}{V_0}} \mathfrak{R}_1 = \mathfrak{L}_1$ ganz im Körper $\sqrt[3]{\frac{1}{V_3}} \mathfrak{R}_2 = \mathfrak{L}_2$ enthalten; infolgedessen gilt:

$$V(\mathfrak{L}_2, \mathfrak{L}_2, \mathfrak{L}_2) \geq V(\mathfrak{L}_1, \mathfrak{L}_1, \mathfrak{L}_2),$$

d. h.

$$1 \geq \frac{d^2 V_1}{\sqrt[3]{V_0^2 V_3}}, \quad \frac{1}{d^2} - 1 \geq \frac{V_1}{\sqrt[3]{V_0^2 V_3}} - 1.$$

Verbinden wir hiermit die Ungleichung (76), so folgt

$$(78) \quad \frac{1}{d^2} - 1 \geq \frac{1}{2^{10} \cdot 3^4 \cdot 7^{\frac{4}{3}}} \frac{(D-1)^6}{D^5}.$$

Vertauschen wir die Rollen von $\mathfrak{R}_1$ und $\mathfrak{R}_2$, so ist D durch $\frac{1}{d}$ und d durch $\frac{1}{D}$ zu ersetzen und dürfen wir mithin in dieser Relation (78) noch D und $\frac{1}{d}$ miteinander vertauschen.

33. Nehmen wir für $\mathfrak{R}_2$ eine Kugel vom Radius 1, so ist $V_3 = \frac{4\pi}{3}$ und $3 V_1$ definiert uns die Oberfläche von $\mathfrak{R}_1$. Die aus (62) entstehende Ungleichung

$$\sqrt[2]{\frac{3 V_1}{4 \pi}} \geq \sqrt[3]{\frac{V_0}{\frac{4 \pi}{3}}}$$

besagt:

Jeder konvexe Körper, welcher nicht eine Kugel ist, besitzt immer eine größere Oberfläche als eine Kugel von demselben Volumen.)*

Es sei jetzt $\mathfrak{K}_1$ zunächst ein vollkommenes Ovaloid. Wir bezeichnen mit df ein Element der Begrenzung von $\mathfrak{K}_1$, mit α, β, γ die äußere Normale dieses Elements, mit $d\omega^*$ ein Element der Kugelfläche $\mathfrak{G}$, mit $\alpha^*, \beta^*, \gamma^*$ die äußere Normale von $d\omega^*$. Alsdann hat die orthogonale Projektion der Begrenzung von $\mathfrak{K}_1$ auf eine zur Richtung $\alpha^*, \beta^*, \gamma^*$ senkrechte Ebene einen Flächeninhalt

$$P(\alpha^*, \beta^*, \gamma^*) = \frac{1}{2} \int |\alpha \alpha^* + \beta \beta^* + \gamma \gamma^*| df.$$

Das arithmetische Mittel aller Größen $P(\alpha^*, \beta^*, \gamma^*)$ in bezug auf alle möglichen Richtungen $(\alpha^*, \beta^*, \gamma^*)$ wird nun

$$\frac{\int P(\alpha^*, \beta^*, \gamma^*) d\omega^*}{\int d\omega^*} = \frac{1}{8\pi} \iint |\alpha \alpha^* + \beta \beta^* + \gamma \gamma^*| df d\omega^* = \frac{1}{4} \int df = \frac{3}{4} V_1,$$

d. i. gleich dem vierten Teile der Oberfläche von $\mathfrak{K}_1$. Dieses Resultat überträgt sich sofort von vollkommenen Ovaloiden auf beliebige konvexe Körper. Verbinden wir damit den vorhin gefundenen Satz, so ergibt sich die Folgerung:

Jeder konvexe Körper, der nicht eine Kugel ist, besitzt stets unter den ihm umbeschriebenen Zylindern solche, welche einen größeren Querschnitt haben als die einer Kugel von demselben Volumen wie der Körper umbeschriebenen Zylinder.

34. Wir nehmen ferner für $\mathfrak{K}_2$ den Würfel von der Kante 1:

$$-\frac{1}{2} \leq x \leq \frac{1}{2}, \quad -\frac{1}{2} \leq y \leq \frac{1}{2}, \quad -\frac{1}{2} \leq z \leq \frac{1}{2};$$

für diesen Körper ist $H_2(\alpha, \beta, \gamma) = \frac{1}{2}(|\alpha| + |\beta| + |\gamma|)$ und stellt infolgedessen $3V_1 = 3V(\mathfrak{K}_2, \mathfrak{K}_1, \mathfrak{K}_1)$ nach (24) und (18) die Summe der Flächeninhalte der drei senkrechten Projektionen des Körpers $\mathfrak{K}_1$ auf die drei Koordinatenebenen dar, während $V_3 = 1$ ist. Aus unserer allgemeinen Ungleichung (62) folgt hier der Satz:

Wird ein konvexer Körper vom Volumen V senkrecht auf drei zueinander orthogonale Ebenen projiziert, so ist die Summe der Flächeninhalte dieser drei Projektionen stets $\geq 3V^{\frac{2}{3}}$ und nur dann $= 3V^{\frac{2}{3}}$, wenn der Körper ein Würfel mit Seitenflächen parallel den drei Ebenen ist. Unter solchen drei zueinander orthogonalen Projektionen des Körpers befindet sich dann gewiß wenigstens eine von einem Flächeninhalt $\geq V^{\frac{2}{3}}$.

*) Von der Literatur über diesen Satz seien erwähnt: Steiner, Über Maximum und Minimum bei den Figuren usw., Crelles Journal, Bd. 24, 1842 (auch Werke, Bd. II). — H. A. Schwarz, Göttinger Nachrichten, 1884 (auch Ges. Abhandlungen, Bd. II). — J. O. Müller, Inauguraldissertation, Göttingen, 1903.

§ 7. Weitere Ungleichungen für die gemischten Volumina.

35. Es seien wieder $\mathfrak{R}_1$ und $\mathfrak{R}_2$ zwei beliebige konvexe Körper und es mögen für sie V_0, V_1, V_2, V_3 die bisherige Bedeutung (s. (61)) haben. Das Volumen eines Körpers $s_1\mathfrak{R}_1 + s_2\mathfrak{R}_2$, wobei $s_1 > 0$ und $s_2 > 0$ ist, hat den Ausdruck

$$s_1^3 V_0 + 3s_1^2 s_2 V_1 + 3s_1 s_2^2 V_2 + s_2^3 V_3.$$

Wenden wir diese Formel auf einen Körper

$$(1+t)\mathfrak{R}_1 + ts\mathfrak{R}_2 = \mathfrak{R}_1 + t(\mathfrak{R}_1 + s\mathfrak{R}_2)$$

an, wo t und $s > 0$ seien, und ordnen die entstehende Formel nach t , so kommen wir zu folgender Bemerkung:

Wenn $\mathfrak{R}_1, \mathfrak{R}_2$ durch $\mathfrak{R}_1, \mathfrak{R}_1 + s\mathfrak{R}_2$ ersetzt werden, so treten an die Stellen von V_0, V_1, V_2, V_3 die Größen

$$V_0, V_0 + sV_1, V_0 + 2sV_1 + s^2V_2, V_0 + 3sV_1 + 3s^2V_2 + s^3V_3.$$

Aus der Ungleichung $V_1^3 - V_0^2V_3 \geq 0$ entsteht nun bei dieser Substitution die Ungleichung

$$\begin{aligned} & (V_0 + sV_1)^3 - V_0^2(V_0 + 3sV_1 + 3s^2V_2 + s^3V_3) \\ & = s^2(3V_0(V_1^2 - V_0V_2) + s(V_1^3 - V_0^2V_3)) \geq 0. \end{aligned}$$

Lassen wir hierin die positive Größe s nach Null abnehmen, so gewinnen wir die folgende in den Größen V_i quadratische Ungleichung:

$$(79) \quad V_1^2 \geq V_0V_2.$$

36. Vertauschen wir die Rollen der Körper $\mathfrak{R}_1$ und $\mathfrak{R}_2$, so treten an die Stelle von V_0, V_1, V_2, V_3 diese Größen in umgekehrter Folge und aus (79) geht die andere Ungleichung

$$(80) \quad V_2^2 \geq V_1V_3$$

hervor.

Andererseits führen die zwei quadratischen Ungleichungen (79) und (80) bei Elimination von V_2 zu

$$V_1^3 \geq \frac{V_0^2V_2^2}{V_1} \geq V_0^2V_3,$$

mithin zurück zu der kubischen Ungleichung, von der wir ausgegangen sind, so daß jene kubische Ungleichung und die quadratische (79) einander gegenseitig zur Folge haben.

Doch besteht in der Abhängigkeit dieser Ungleichungen voneinander ein wesentlicher Unterschied, insofern als die Grenzfälle der quadratischen Ungleichung sofort auch die Grenzfälle der kubischen ergeben, dagegen nicht umgekehrt aus den Bedingungen, unter welchen in der kubischen Ungleichung das Gleichheitszeichen statthat, vollkommen zu ersehen ist, wann das Gleichheitszeichen in der quadratischen eintritt.

37. Man kann nun die quadratische Ungleichung $V_1^2 \geq V_0 V_2$ auch auf einem direkten Wege durch analoge Überlegungen, wie sie zu der genaueren Ungleichung (76) führten, herleiten und gelangt alsdann auch zur Kenntnis der Grenzfälle dieser quadratischen Ungleichung. Wir begnügen uns hier, das bezügliche Resultat ohne Beweis anzugeben.

Wir bezeichnen eine Stützebene $\varphi \leq 0$ an einen konvexen Körper $\mathfrak{K}$ als eine *Eckstützebene* an $\mathfrak{K}$, wenn wir $\varphi = t_1 \varphi_1 + t_2 \varphi_2 + t_3 \varphi_3$ setzen können, so daß t_1, t_2, t_3 sämtlich > 0 und $\varphi_1 \leq 0, \varphi_2 \leq 0, \varphi_3 \leq 0$ drei solche Stützebenen an $\mathfrak{K}$ sind, deren äußere Normalenrichtungen *unabhängig* sind, d. h. nicht einer einzigen Ebene angehören. Ein konvexer Körper $\mathfrak{K}$ heiße ein *Kappenkörper* eines konvexen Körpers $\mathfrak{K}'$, wenn mit Ausnahme nur von gewissen Eckstützebenen alle anderen Stützebenen an $\mathfrak{K}$ zugleich auch Stützebenen an $\mathfrak{K}'$ bilden. Der allgemeinste *Kappenkörper einer Kugel* wird erhalten, indem man die Kugel als Grundbestandteil nimmt, auf ihrer Oberfläche eine endliche oder unendliche Anzahl von Kalotten bezeichnet, von denen keine die Größe einer halben Kugelfläche erreicht oder überschreitet und die untereinander höchstens in den Randpunkten zusammenstoßen, und sodann die Kugel auf jeder dieser Kalotten mit der Kegelkappe versieht, bei der die Kalotte als Basis dient und die Erzeugenden des Mantels die Kugel in dem Rande der Kalotte berühren.

Es besteht nun der Satz:

In der Ungleichung $V_1^2 \geq V_0 V_2$ für zwei beliebige konvexe Körper $\mathfrak{K}_1$ und $\mathfrak{K}_2$ hat dann und nur dann das Gleichheitszeichen statt, wenn $\mathfrak{K}_1$ homothetisch mit $\mathfrak{K}_2$ oder mit einem Kappenkörper von $\mathfrak{K}_2$ ist.

Bei einem vollkommenen Ovaloide kommen keine Eckstützebenen vor und kann ein solches daher niemals Kappenkörper eines anderen konvexen Körpers sein. Wenn also $\mathfrak{K}_1$ ein vollkommenes Ovaloid ist, gilt in $V_1^2 \geq V_0 V_2$ das Gleichheitszeichen nur dann, wenn $\mathfrak{K}_2$ und $\mathfrak{K}_1$ selbst homothetisch sind.

Nehmen wir für $\mathfrak{K}_2$ eine Kugel vom Radius 1, so ist V_0 das Volumen, $3 V_1$ die Oberfläche, $3 V_2$ das Integral der mittleren Krümmung von $\mathfrak{K}_1$ und $V_3 = \frac{4\pi}{3}$. Als dann gilt $V_1^2 \geq V_0 V_2$ und hat hier das Gleichheitszeichen dann und nur dann statt, wenn $\mathfrak{K}_1$ eine Kugel oder ein Kappenkörper einer Kugel ist. Zweitens gilt $V_2^2 \geq V_1 V_3$ und in dieser Ungleichung hat das Gleichheitszeichen dann und nur dann statt, wenn $\mathfrak{K}_1$ selbst eine Kugel ist.

Danach liefern unter allen konvexen Körpern von gleicher Oberfläche die Kugeln und die Kappenkörper von Kugeln das Maximum des Produkts aus Volumen und Integral der mittleren Krümmung und andererseits allein die Kugeln das Minimum des Integrals der mittleren Krümmung. Aus

diesen beiden Tatsachen zusammen folgt, daß unter allen konvexen Körpern von gleicher Oberfläche die Kugeln das größte Volumen darbieten.

38. Es seien jetzt $\mathfrak{K}_1, \mathfrak{K}_2, \mathfrak{K}_3$ drei beliebige konvexe Körper; das Volumen eines Körpers $\mathfrak{K} = s_1 \mathfrak{K}_1 + s_2 \mathfrak{K}_2 + s_3 \mathfrak{K}_3$, wobei $s_1, s_2, s_3 \geq 0$, aber nicht durchweg 0 sind, stellt sich durch eine ternäre kubische Form

$$V = \sum V_{jkl} s_j s_k s_l$$

dar, wo jeder der Indizes die Werte 1, 2, 3 zu durchlaufen hat. Der Koeffizient V_{jkl} bedeutet das gemischte Volumen $V(\mathfrak{K}_j, \mathfrak{K}_k, \mathfrak{K}_l)$.

Wir wenden nun die für zwei konvexe Körper nachgewiesene Ungleichung $V_1^2 - V_0 V_2 \geq 0$ derart an, daß wir für den ersten Körper $p_1 \mathfrak{K}_1 + p_2 \mathfrak{K}_2 + p_3 \mathfrak{K}_3$, für den zweiten $q_1 \mathfrak{K}_1 + q_2 \mathfrak{K}_2 + q_3 \mathfrak{K}_3$ nehmen, wobei $p_1, \dots, q_3$ lauter positive Parameter seien. Setzen wir

$$V_{jk1} p_1 + V_{jk2} p_2 + V_{jk3} p_3 = P_{jk},$$

so wird hier

$$V_0 = \sum P_{jk} p_j p_k, \quad V_1 = \sum P_{jk} p_j q_k, \quad V_2 = \sum P_{jk} q_j q_k.$$

Wir nehmen noch

$$q_1 = t p_1 + r_1, \quad q_2 = t p_2 + r_2, \quad q_3 = t p_3 + r_3$$

an; dabei können r_1, r_2, r_3 beliebige reelle Werte bedeuten und bei positiven p_1, p_2, p_3 kann stets t so groß gewählt werden, daß auch q_1, q_2, q_3 sämtlich > 0 sind. Die in Rede stehende Ungleichung verwandelt sich nunmehr in

$$\left(\sum P_{jk} p_j r_k \right)^2 - \left(\sum P_{jk} p_j p_k \right) \left(\sum P_{jk} r_j r_k \right) \geq 0.$$

Wir nehmen jetzt $p_3 = 1$ an und lassen die positiven Werte p_1, p_2 sich der Grenze Null nähern; es entsteht dann hieraus:

$$(V_{133} r_1 + V_{233} r_2)^2 - V_{333} (V_{113} r_1^2 + 2 V_{123} r_1 r_2 + V_{223} r_2^2) \geq 0$$

und dabei muß diese Ungleichung für beliebige Werte r_1, r_2 gelten. Setzen wir nun

$$\Delta^{(3)} = \begin{vmatrix} V_{113} & V_{123} & V_{133} \\ V_{213} & V_{223} & V_{233} \\ V_{313} & V_{323} & V_{333} \end{vmatrix}$$

und bezeichnen die adjungierte Unterdeterminante zum Element V_{jks} hier mit $W_{jk}^{(3)}$, so schreibt sich diese Ungleichung

$$(81) \quad -W_{22}^{(3)} r_1^2 + 2 W_{21}^{(3)} r_1 r_2 - W_{11}^{(3)} r_2^2 \geq 0.$$

Die Determinante der binären quadratischen Form rechts hier ist $V_{333} \Delta^{(3)}$ und folgt daher

$$\Delta^{(3)} \geq 0.$$

Ferner haben wir insbesondere $-W_{22}^{(3)} \geq 0$, $-W_{11}^{(3)} \geq 0$. Ist nun etwa

— $W_{22}^{(3)} > 0$, so geht vermöge der Beziehung

$$W_{22}^{(3)} W_{33}^{(3)} - (W_{23}^{(3)})^2 = V_{113} \Delta^{(3)}$$

die Ungleichung

$$-W_{33}^{(3)} = V_{123}^2 - V_{113} V_{223} \geq 0$$

hervor. Analog erschließt man diese Ungleichung, wenn $-W_{11}^{(3)} > 0$ ist.

Hat man jedoch sowohl $-W_{22}^{(3)} = 0$ wie $-W_{11}^{(3)} = 0$, so muß, da die Ungleichung (81) für beliebige r_1 und r_2 statthaben soll, durchaus auch $W_{12}^{(3)} = 0$ sein, und dann sind in $\Delta^{(3)}$ die erste und dritte und ferner die zweite und dritte Reihe proportional und folgt dadurch

$$\Delta^{(3)} = 0, \quad -W_{33}^{(3)} = 0.$$

In allen Fällen gilt hiernach

$$(82) \quad V_{123}^2 \geq V_{113} V_{223}.$$

Verbinden wir diese Ungleichung mit den zwei Ungleichungen

$$(83) \quad V_{113}^3 \geq V_{111}^2 V_{333}, \quad V_{223}^3 \geq V_{222}^2 V_{333},$$

welche die Beziehung $V_1^3 \geq V_0^2 V_3$ für $\mathfrak{R}_1$ und $\mathfrak{R}_3$ bzw. für $\mathfrak{R}_2$ und $\mathfrak{R}_3$ vorstellen, so gelangen wir durch Elimination von V_{113} und V_{223} zu

$$(84) \quad V_{123}^3 \geq V_{111} V_{222} V_{333}.$$

Dabei kann in dieser Ungleichung das Gleichheitszeichen nur statthaben, wenn in *beiden* Ungleichungen (83) das Gleichheitszeichen gilt. Damit kommen wir zu folgendem Satze:

Drei beliebige konvexe Körper vom Volumen 1 (hier $\frac{1}{\sqrt[3]{V_{111}}} \mathfrak{R}_1$, $\frac{1}{\sqrt[3]{V_{222}}} \mathfrak{R}_2$, $\frac{1}{\sqrt[3]{V_{333}}} \mathfrak{R}_3$) ergeben stets ein gemischtes Volumen (nämlich $\frac{V_{123}}{\sqrt[3]{V_{111} V_{222} V_{333}}}$), das ≥ 1 und nur dann $= 1$ ist, wenn alle drei Körper einander homothetisch sind.

Die frühere Ungleichung $V_1^3 \geq V_0^2 V_3$ geht aus diesem Satze, ebenso die Ungleichung $V_1^2 \geq V_0 V_2$ aus (82) als spezieller Fall hervor, wenn der zweite und dritte Körper identisch gesetzt werden.

39. Es seien $\mathfrak{R}_1, \mathfrak{R}_2, \dots, \mathfrak{R}_m$ m beliebige konvexe Körper. Wir wenden die Ungleichung (82) für drei konvexe Körper an, indem wir für den ersten einen Körper $\sum p_i \mathfrak{R}_i$, für den zweiten einen Körper $\sum (tp_i + r_i) \mathfrak{R}_i$, für den dritten einen Körper $\sum s_i \mathfrak{R}_i$ ($i = 1, 2, \dots, m$) nehmen, wobei die r_i beliebige Größen und die Werte $p_i, tp_i + r_i, s_i$ sämtlich positiv sind. Setzen wir

$$V(\mathfrak{R}_j, \mathfrak{R}_k, \mathfrak{R}_l) = V_{jkl}, \quad V_{jk1}s_1 + V_{jk2}s_2 + \dots + V_{jkm}s_m = S_{jk},$$

so wird die betreffende Ungleichung

$$\left(\sum S_{jk} p_j r_k \right)^2 - \left(\sum S_{jk} p_j p_k \right) \left(\sum S_{jk} r_j r_k \right) \geq 0;$$

sie bedeutet, daß die quadratische Form $\sum S_{jk} r_j r_k$ bei einer Transformation in ein Aggregat von Quadraten reeller unabhängiger linearer Formen mit positiven oder negativen Vorzeichen *ein einziges Quadrat mit positivem und im übrigen lauter Quadrate mit negativem Vorzeichen* darbieten muß. Im besonderen muß danach die Determinante dieser Form mit $(-1)^{m-1}$ multipliziert, stets ≥ 0 sein, d. h. für beliebige positive Werte $s_1, s_2, \dots, s_m$ muß stets die Ungleichung

$$(85) \quad (-1)^{m-1} |V_{jk1}s_1 + V_{jk2}s_2 + \dots + V_{jkm}s_m| \geq 0$$

bestehen.

§ 8. Stetig gekrümmte konvexe Körper.

40. Das allgemeine Theorem $V_1^3 \geq V_0^2 V_3$ für zwei beliebige konvexe Körper ist aufs engste mit gewissen Fragen über die Krümmung konvexer Körper verknüpft.

Sind $\mathfrak{R}$ und $\mathfrak{R}'$ zwei beliebige konvexe Körper, so wollen wir den Wert $3V(\mathfrak{R}', \mathfrak{R}, \mathfrak{R})$ die *Relativoberfläche von $\mathfrak{R}$ in bezug auf $\mathfrak{R}'$* nennen. Es seien H und H' die Stützebenenfunktionen von $\mathfrak{R}$ und $\mathfrak{R}'$. Wir nehmen $\mathfrak{R}$ zunächst als ein vollkommenes Ovaloid an und verwenden dafür die in § 3 eingeführten Bezeichnungen; alsdann gilt nach (24) und (18) der Ausdruck

$$(86) \quad 3V(\mathfrak{R}', \mathfrak{R}, \mathfrak{R}) = \int H'(RT - S^2) d\omega = \int H' df;$$

darin bedeutet df das Flächenelement der Begrenzung von $\mathfrak{R}$, welches durch parallele Normalen dem Element $d\omega$ der Kugelfläche $\mathfrak{E}$ entspricht, und $RT - S^2$ das Produkt der Hauptkrümmungsradien, die reziproke Gaußsche Krümmung, an der Stelle von df .

Setzen wir $H' = H + \delta H$, so hat der Körper $(1-t)\mathfrak{R} + t\mathfrak{R}'$, wenn $t > 0$ und < 1 ist, die Stützebenenfunktion $H + t\delta H$. Die Differenz aus dem Volumen dieses Körpers und dem Volumen von $\mathfrak{R}$ ist dann bis auf Größen von der Ordnung t^2 gleich

$$t \int \delta H(RT - S^2) d\omega = t \int \delta H df.$$

41. Diese Beziehungen, die jedenfalls bei einem vollkommenen Ovaloide $\mathfrak{R}$ statthaben, veranlassen uns nun zu folgender Definition:

Ein konvexer Körper $\mathfrak{R}$ soll stetig gekrümmt heißen, wenn für ihn eine auf der Kugeloberfläche $\mathfrak{E}$ stetige und durchweg positive Funktion $F = F(\alpha, \beta, \gamma)$ existiert derart, daß die Relativoberfläche von $\mathfrak{R}$ in bezug auf einen beliebigen konvexen Körper $\mathfrak{R}'$ mit der Stützebenenfunktion H' immer den Ausdruck hat:

$$(87) \quad 3V(\mathfrak{R}', \mathfrak{R}, \mathfrak{R}) = \int H' F d\omega.$$

Diese Funktion $F(\alpha, \beta, \gamma)$ (oder $F(\vartheta, \psi)$) wollen wir dann die *Krümmungsfunktion* des Körpers $\mathfrak{R}$ nennen. Zunächst leuchtet ein, daß bei einem stetig gekrümmten Körper $\mathfrak{R}$ diese Funktion jedenfalls eine durchaus bestimmte ist. Denn nehmen wir an, es gäbe für $\mathfrak{R}$ noch eine zweite auf $\mathfrak{E}$ stetige und positive Funktion $F^*(\alpha, \beta, \gamma)$, welche in derselben Weise wie F in (87) eintreten könnte, so wäre dann für jeden beliebigen konvexen Körper $\mathfrak{R}'$ stets

$$(88) \quad \int H' F^* d\omega = \int H' F d\omega, \quad \int H' (F^* - F) d\omega = 0.$$

Da $F^* - F$ auf $\mathfrak{E}$ nicht durchweg Null ist, können wir auf $\mathfrak{E}$ einen Punkt finden, wo $F^* - F \neq 0$ ist, und hernach können wir, da $F^* - F$ auf $\mathfrak{E}$ stetig ist, eine ganze Kalotte um diesen Punkt bestimmen, (die wir kleiner als eine Halbkugel wählen), so daß $F^* - F$ daselbst überall $\neq 0$ und also von konstantem Vorzeichen, etwa stets > 0 ist. Wir setzen auf die Kalotte die Kegelkappe, bei welcher die Kalotte die Basis bildet und die Erzeugenden des Mantels die Kugelfläche $\mathfrak{E}$ im Rande der Kalotte berühren, und nehmen nun in der Gleichung (88) für $\mathfrak{R}'$ einmal den Körper, der aus der Kugel $\mathfrak{G}$ und dieser Kappe besteht, ein andres Mal $\mathfrak{G}$ selbst, so hat das Integral $\int H' (F^* - F) d\omega$ im ersten Falle einen größeren Wert als im zweiten Falle und kann daher nicht in beiden Fällen Null sein.

42. Nehmen wir mit dem Körper $\mathfrak{R}'$, für den (87) gebildet ist, die Translation vom Nullpunkte nach einem Punkte a, b, c vor, so ist $H'(\alpha, \beta, \gamma)$ durch

$$H'(\alpha, \beta, \gamma) + a\alpha + b\beta + c\gamma$$

zu ersetzen. Da nun hierbei der Wert $3V(\mathfrak{R}', \mathfrak{R}, \mathfrak{R})$ sich nicht ändert und die Größen a, b, c völlig beliebig sind, so ersehen wir:

Die Krümmungsfunktion F für einen stetig gekrümmten Körper muß stets die drei Bedingungsgleichungen erfüllen:

$$(89) \quad \int \alpha F d\omega = 0, \quad \int \beta F d\omega = 0, \quad \int \gamma F d\omega = 0.$$

Andererseits erhellt, daß, wenn wir mit $\mathfrak{R}$ eine Translation vornehmen, dabei die Krümmungsfunktion invariant ist. Unter den sämtlichen, aus $\mathfrak{R}$ durch Translationen abzuleitenden Körpern können wir einen bestimmten Körper durch die Lage des Schwerpunkts fixieren.

Ist t ein positiver Parameter, so hat $t\mathfrak{R}$ die Krümmungsfunktion $t^2 F$, wenn F die Krümmungsfunktion von $\mathfrak{R}$ ist; dieses folgt unmittelbar aus der Formel

$$3V(\mathfrak{R}', t\mathfrak{R}, t\mathfrak{R}) = 3t^2 V(\mathfrak{R}', \mathfrak{R}, \mathfrak{R}).$$

43. Unsere weiteren Entwicklungen nun werden darauf gerichtet sein, den folgenden sehr bemerkenswerten Satz zu beweisen:

Ist $F(\alpha, \beta, \gamma)$ eine auf der Kugelfläche $\mathfrak{E}$ beliebig vorgeschriebene, daselbst stetige und durchweg positive Funktion, welche die drei Bedingungsgleichungen

$$\int \alpha F d\omega = 0, \quad \int \beta F d\omega = 0, \quad \int \gamma F d\omega = 0$$

erfüllt, so gibt es immer einen stetig gekrümmten konvexen Körper $\mathfrak{K}$ mit $F(\alpha, \beta, \gamma)$ als Krümmungsfunktion, und dieser Körper ist bestimmt bis auf eine willkürliche Translation, durch die man ihn noch abändern kann, also eindeutig festgelegt, wenn noch zusätzlich gefordert wird, daß sein Schwerpunkt im Nullpunkte liegen soll.

Wir beweisen zunächst den Nachsatz, daß den hier gestellten Forderungen, wenn überhaupt, jedenfalls nur durch einen einzigen Körper genügt werden kann. In der Tat, nehmen wir an, es sei ein konvexer Körper $\mathfrak{K}$ gefunden mit F als Krümmungsfunktion und mit dem Schwerpunkte im Nullpunkte, und es sei H die Stützebenenfunktion von $\mathfrak{K}$. Zunächst ist klar, daß unter den mit $\mathfrak{K}$ homothetischen Körpern kein anderer diese beiden Eigenschaften mit $\mathfrak{K}$ teilt.

Das Volumen von $\mathfrak{K}$ ist durch

$$V = \frac{1}{3} \int H F d\omega$$

dargestellt. Ist sodann $\mathfrak{K}'$ mit der Stützebenenfunktion H' und dem Volumen V' ein beliebiger konvexer Körper, der mit $\mathfrak{K}$ nicht homothetisch ist, so ergibt das allgemeine Theorem $\sqrt[3]{\frac{V_1}{V_s}} > \sqrt[3]{\frac{V_0}{V_0}}$ auf die Körper $\mathfrak{K}$ und $\mathfrak{K}'$ angewandt:

$$(90) \quad \frac{1}{3} \int \frac{H'}{\sqrt[3]{V'}} F d\omega > \frac{1}{3} \int \frac{H}{\sqrt[3]{V}} F d\omega.$$

Darin sind nun $\frac{H}{\sqrt[3]{V}}$ und $\frac{H'}{\sqrt[3]{V'}}$ die Stützebenenfunktionen der Körper

$$\mathfrak{L} = \frac{1}{\sqrt[3]{V}} \mathfrak{K} \quad \text{und} \quad \mathfrak{L}' = \frac{1}{\sqrt[3]{V'}} \mathfrak{K}',$$

und diese zwei mit $\mathfrak{K}$ bzw. $\mathfrak{K}'$ homothetischen Körper sind dadurch charakterisiert, daß sie ein Volumen = 1 haben.

Damit kommen wir zu folgendem Ergebnisse, welches zeigt, daß der Körper $\mathfrak{K}$ eindeutig bestimmt ist.

Man betrachte für alle möglichen konvexen Körper $\mathfrak{L}$ vom Volumen 1 das Integral

$$J(\mathfrak{L}) = \frac{1}{3} \int L F d\omega,$$

wobei $L = L(\alpha, \beta, \gamma)$ die Stützebenenfunktion von $\mathfrak{L}$ und $F = F(\alpha, \beta, \gamma)$ die gegebene Funktion für die Argumente α, β, γ , die Normale in $d\omega$, be-

deute; gibt es einen stetig gekrümmten konvexen Körper $\mathfrak{K}$ mit F als Krümmungsfunktion, so hat dieses Integral $J(\mathfrak{L})$ für einen ganz bestimmten Körper $\mathfrak{L}$ mit dem Nullpunkte als Schwerpunkt (und für die aus ihm durch Translationen hervorgehenden Körper) einen kleinsten Wert J , und alsdann ist $\mathfrak{K}$ identisch mit $\sqrt{J}\mathfrak{L}$ oder geht aus letzterem Körper durch eine Translation hervor.

44. Zugleich werden wir hierdurch auf ein gewisses Variationsproblem hingewiesen, das in nahem Zusammenhange mit der von uns zu behandelnden Aufgabe steht; und in der Tat erkennen wir sofort, daß die Differentialgleichung zu diesem Variationsproblem

$$RT - S^2 = F(\vartheta, \psi)$$

wird, worin F die gegebene, den Gleichungen (89) genügende Funktion ist, und die in R, S, T auftretende Funktion $H(\vartheta, \psi)$ gefunden werden soll. Diese Gleichung ist nach den Ausdrücken von R, S, T in (17) für $H(\vartheta, \psi)$ eine quadratische Differentialgleichung zweiter Ordnung von dem Monge-Ampèrèschen Typus; sie transformiert sich, wenn wir auf der Oberfläche des gesuchten Körpers die Koordinate z als Funktion von x, y einführen, in die Gleichung:

$$rt - s^2 = f(p, q)$$

für diese Funktion $z(x, y)$, wobei

$$\frac{\partial z}{\partial x} = p, \quad \frac{\partial z}{\partial y} = q, \quad \frac{\partial^2 z}{\partial x^2} = r, \quad \frac{\partial^2 z}{\partial x \partial y} = s, \quad \frac{\partial^2 z}{\partial y^2} = t$$

gesetzt ist und

$$\frac{f(p, q)}{(1 + p^2 + q^2)^2} = F$$

als eine beliebige eindeutige und stetige Funktion in

$$\alpha = \frac{-p}{\sqrt{1 + p^2 + q^2}}, \quad \beta = \frac{-q}{\sqrt{1 + p^2 + q^2}}, \quad \gamma = \frac{1}{\sqrt{1 + p^2 + q^2}}$$

vorgeschrieben sein kann, die nur den Gleichungen (89) zu genügen hat.

Wir werden jetzt ein gewisses Problem mit einer nur *endlichen* Anzahl von Parametern, sozusagen eine Art von Differenzengleichung erledigen und hernach von diesem Probleme aus durch einen Grenzübergang zur Lösung der hier gestellten Aufgabe, die von einer kontinuierlichen Funktion F handelt, gelangen.

§ 9. Bestimmung eines Polyeders durch die Normalen und die Inhalte der Seitenflächen.

45. Es sei $\mathfrak{P}$ ein beliebiges Polyeder mit n Seitenflächen, die in irgendeiner Ordnung numeriert seien. Wir wollen annehmen, der Nullpunkt liege in $\mathfrak{P}$ selbst, sei es im Inneren, sei es auf der Begrenzung

von $\mathfrak{P}$. Wir bezeichnen für $i = 1, 2, \dots, n$ mit $(\alpha_i, \beta_i, \gamma_i)$ die äußere Normale, mit F_i den Flächeninhalt der i^{ten} Seitenfläche von $\mathfrak{P}$, mit p_i die Länge des vom Nullpunkt auf die Ebene dieser Fläche gefällten Lotes, endlich mit V das Volumen von $\mathfrak{P}$.

Die Richtungen $(\alpha_i, \beta_i, \gamma_i)$ sind gewiß so beschaffen, daß sie nicht sämtlich einer Ebene angehören; die Größen F_i sind sämtlich > 0 . Bei einer Translation des Polyeders $\mathfrak{P}$ bleiben alle Größen $\alpha_i, \beta_i, \gamma_i, F_i$ ungeändert.

Indem wir das Polyeder $\mathfrak{P}$ nach der in § 2 entwickelten Methode als Grenze von vollkommenen Ovaloiden darstellen und die Formel (86) heranziehen, gewinnen wir die folgende Regel:

Ist $\mathfrak{Q}$ ein beliebiger konvexer Körper, Q seine Stützebenenfunktion und setzen wir allgemein $Q(\alpha_i, \beta_i, \gamma_i) = q_i$, so wird das gemischte Volumen

$$(91) \quad V(\mathfrak{Q}, \mathfrak{P}, \mathfrak{P}) = \frac{1}{3} (F_1 q_1 + F_2 q_2 + \dots + F_n q_n).$$

Insbesondere entnehmen wir daraus für das Volumen von $\mathfrak{P}$ den Ausdruck

$$(92) \quad V = \frac{1}{3} (F_1 p_1 + F_2 p_2 + \dots + F_n p_n).$$

Genau so, wie wir aus (87) die drei Gleichungen (89) folgerten, schließen wir aus (91) durch beliebige Translation des Körpers $\mathfrak{Q}$ auf das notwendige Bestehen der drei Gleichungen:

$$(93) \quad \sum F_i \alpha_i = 0, \quad \sum F_i \beta_i = 0, \quad \sum F_i \gamma_i = 0 \quad (i = 1, 2, \dots, n).$$

Bezeichnen wir das Volumen von $\mathfrak{Q}$ mit V_* und bilden wir unsere allgemeine Ungleichung $\frac{{}^3 V_1}{{}^3 \sqrt{V_*}} \geq \frac{{}^3 V_0}{{}^3 \sqrt{V_0}}$ in bezug auf das Polyeder $\mathfrak{P}$ als ersten und den Körper $\mathfrak{Q}$ als zweiten Körper, so finden wir, daß stets

$$(94) \quad \frac{F_1 q_1 + F_2 q_2 + \dots + F_n q_n}{V_*^{\frac{1}{3}}} \geq \frac{F_1 p_1 + F_2 p_2 + \dots + F_n p_n}{V^{\frac{1}{3}}}$$

ist und hierin das Gleichheitszeichen nur dann gilt, wenn $\mathfrak{Q}$ mit $\mathfrak{P}$ homothetisch ist.

Wir werden nunmehr den folgenden Satz beweisen*):

Es seien n beliebige Richtungen $(\alpha_i, \beta_i, \gamma_i)$ für $i = 1, 2, \dots, n$ gegeben, die nicht sämtlich einer Ebene angehören, und dazu n beliebige positive Größen F_i , so daß die drei Gleichungen bestehen

$$\sum F_i \alpha_i = 0, \quad \sum F_i \beta_i = 0, \quad \sum F_i \gamma_i = 0 \quad (i = 1, 2, \dots, n);$$

alsdann gibt es stets ein Polyeder $\mathfrak{P}$ mit n Seitenflächen, wofür die Richtungen

*) Diesen Satz habe ich zuerst in dem Aufsätze: *Allgemeine Lehrsätze über die konvexen Polyeder*, Göttinger Nachrichten, 1897, publiziert. (Diese Ges. Abhandlungen, Bd. II, S. 103.

$(\alpha_i, \beta_i, \gamma_i)$ die äußeren Normalen und die Größen F_i die Flächeninhalte dieser Seitenflächen bilden, und dieses Polyeder $\mathfrak{P}$ ist vollkommen bestimmt bis auf eine beliebige Translation, durch die man dasselbe noch variieren kann, also eindeutig festgelegt, wenn ferner noch die Lage seines Schwerpunktes beliebig vorgeschrieben wird.

46. Um zu einem Beweise dieses Satzes zu gelangen, fassen wir jetzt die Gesamtheit aller möglichen Polyeder mit n oder weniger Seitenflächen ins Auge, welche die äußeren Normalen der Seitenflächen nur unter den n Richtungen $(\alpha_i, \beta_i, \gamma_i)$ besitzen und welche ferner den Nullpunkt in sich schließen.

Sind $q_1, q_2, \dots, q_n$ irgendwelche n Größen, sämtlich ≥ 0 und derart, daß die n Ungleichungen

$$(95) \quad \alpha_i x + \beta_i y + \gamma_i z \leq q_i \quad (i = 1, 2, \dots, n)$$

ein wirkliches Polyeder definieren, so bezeichnen wir dieses Polyeder mit $\mathfrak{R}(q_1, q_2, \dots, q_n)$ oder $\mathfrak{R}(q_i)$, sein Volumen mit $V(q_1, q_2, \dots, q_n)$ oder $V(q_i)$. Dabei kann auch ein Teil dieser Ungleichungen eine Folge der übrigen sein, das Polyeder also weniger als n wirkliche Seitenflächen besitzen. Wir bezeichnen ferner mit q_i^* den größten Wert von $\alpha_i x + \beta_i y + \gamma_i z$ in $\mathfrak{R}(q_i)$; für diejenigen Indizes i , wobei $\mathfrak{R}$ eine wirkliche Seitenfläche mit der Normale $(\alpha_i, \beta_i, \gamma_i)$ besitzt, ist immer $q_i^* = q_i$, während wir bei den anderen Indizes nur $q_i \geq q_i^*$ behaupten können. Wir nennen $q_1^*, q_2^*, \dots, q_n^*$ die *tangentialen Parameter* von $\mathfrak{R}(q_1, q_2, \dots, q_n)$; offenbar ist

$$\mathfrak{R}(q_1, q_2, \dots, q_n) = \mathfrak{R}(q_1^*, q_2^*, \dots, q_n^*).$$

Existiert nun ein Polyeder $\mathfrak{P} = \mathfrak{R}(p_1, p_2, \dots, p_n)$ gemäß den Forderungen unseres Satzes, so zeigt die Ungleichung (94), daß unter allen vorhandenen Körpern $\mathfrak{R}(q_1, q_2, \dots, q_n)$ eben dieses Polyeder $\mathfrak{P}$ und die ihm homothetischen Körper den kleinsten Wert des Ausdrucks

$$\frac{F_1 q_1 + F_2 q_2 + \dots + F_n q_n}{(V(q_1, q_2, \dots, q_n))^{\frac{1}{n}}}$$

liefern. Daraus erhellt zunächst der letzte Teil jenes Satzes, daß nämlich das gesuchte Polyeder $\mathfrak{P}$ gewiß nur auf eine Art, abgesehen von einer Translation, bestimmt werden kann.

Nunmehr wollen wir die wirkliche Existenz jenes Polyeders $\mathfrak{P}$ dartun.

47. Wir zeigen vor allem, daß, wie die n Richtungen $(\alpha_i, \beta_i, \gamma_i)$ vorausgesetzt sind, gewiß irgendwelche Polyeder $\mathfrak{R}(q_1, q_2, \dots, q_n)$ vorhanden sind. In der Tat, der durch die n Ungleichungen

$$(96) \quad \alpha_i x + \beta_i y + \gamma_i z \leq 1 \quad (i = 1, 2, \dots, n)$$

definierte Bereich $\mathfrak{R}(1, 1, \dots, 1) = \mathfrak{R}(1)$ stellt ein solches Polyeder, und zwar wirklich mit n Seitenflächen vor. Denn diese Ungleichungen sind

die Bedingungen von n verschiedenen Tangentialebenen an die Kugel $\mathcal{G}$; der Bereich enthält daher die Kugel $\mathcal{G}$ in sich und besitzt jene n Ebenen als extreme Stützebenen. Andererseits kann dieser Bereich $\mathfrak{R}(1)$ sich nicht ins Unendliche erstrecken; denn der Abstand eines beliebigen Punktes x, y, z in ihm von der i^{ten} jener Stützebenen wird

$$t_i = 1 - \alpha_i x - \beta_i y - \gamma_i z \geq 0$$

und entnehmen wir aus den Gleichungen (93) dann

$$\sum F_i t_i = \sum F_i \quad (i = 1, 2, \dots, n).$$

Da die Werte F_i sämtlich > 0 sind, besteht hiernach für eine jede Größe t_i eine obere Grenze. Nun können wir unter den Stützebenen (96) gewiß drei solche herausgreifen, die sich in einem Punkte schneiden; durch die oberen Grenzen der drei zugehörigen Größen t_i werden dann drei zu ihnen parallele Ebenen angewiesen, welche mit den ersteren zusammen ein Parallelepipedum bestimmen, in dem der Bereich $\mathfrak{R}(1)$ ganz enthalten sein muß. Danach liegt dieser Bereich ganz im Endlichen. Das Volumen von $\mathfrak{R}(1)$ setzen wir $= \frac{1}{l^3}$, alsdann hat $\mathfrak{R}(l, l, \dots, l)$ das Volumen 1. Weiter wird überhaupt für beliebige endliche Werte $q_1, q_2, \dots, q_n$ der durch die Ungleichungen (95) bestimmte Bereich ganz im Endlichen liegen und bei lauter positiven Werten q_i stets ein wirkliches Polyeder, wenn auch nicht immer mit n Seitenflächen, vorstellen.

48. Wir betrachten jetzt in der n -fachen Mannigfaltigkeit $\mathfrak{M}$ von n beliebig veränderlichen reellen Größen $q_1, q_2, \dots, q_n$ den Bereich $\mathfrak{B}$ aller solchen Systeme $q_1, q_2, \dots, q_n$, „Punkte“ (q_i), wobei $q_1 \geq 0, q_2 \geq 0, \dots, q_n \geq 0$ ist und den Größen $q_1, q_2, \dots, q_n$ ein Polyeder $\mathfrak{R}(q_1, q_2, \dots, q_n)$ von einem Volumen

$$V(q_1, q_2, \dots, q_n) \geq 1$$

entspricht.

Sind (r_i) und (s_i) ($i = 1, 2, \dots, n$) zwei beliebige Punkte dieses Bereichs $\mathfrak{B}$ und (r_i^*) bzw. (s_i^*) die tangentialen Parameter von $\mathfrak{R}(r_i)$ und $\mathfrak{R}(s_i)$ und ist t ein Wert > 0 und < 1 , so besitzt der aus diesen zwei Polyedern abzuleitende Bereich $(1-t)\mathfrak{R}(r_i) + t\mathfrak{R}(s_i)$ nach (77) ein Volumen

$$\geq ((1-t)\sqrt[3]{V(r_i)} + t\sqrt[3]{V(s_i)})^3 \geq ((1-t) + t)^3 = 1.$$

Die Stützebenenfunktion des letzteren Bereichs hat für die Argumente $\alpha_i, \beta_i, \gamma_i$ den Wert $(1-t)r_i^* + ts_i^*$, und danach ist dieser Bereich entweder identisch, oder, wenn nicht identisch, so doch jedenfalls ganz enthalten in dem Bereich

$$\alpha_i x + \beta_i y + \gamma_i z \leq (1-t)r_i^* + ts_i^* \quad (i = 1, 2, \dots, n),$$

der dadurch gewiß auch ein Polyeder vorstellt, und um so mehr dann

enthalten im Polyeder $\mathfrak{R}((1-t)r_i + ts_i)$. Mithin ist für das letztere Polyeder notwendig ebenfalls das Volumen ≥ 1 , also

$$V((1-t)r_i + ts_i) \geq 1.$$

Der Bereich $\mathfrak{B}$ hat danach die Eigenschaft, mit irgend zwei Punkten $(r_i), (s_i)$ stets die ganze sie verbindende „Strecke“ von Punkten $((1-t)r_i + ts_i)$ zu enthalten, und ist deshalb als ein „konvexes Gebilde“ in der Mannigfaltigkeit $\mathfrak{M}$ anzusprechen. Da $V(q_1, q_2, \dots, q_n)$ eine stetige Funktion der Argumente ist, wird ferner der Bereich $\mathfrak{B}$ abgeschlossen sein und wird seine Begrenzung von denjenigen Punkten (q_i) gebildet werden, für welche

$$V(q_1, q_2, \dots, q_n) = 1$$

ist. Wir haben jetzt nach dem kleinsten Werte des mit den gegebenen positiven Größen F_i zu bildenden linearen Ausdrucks

$$(97) \quad F_1 q_1 + F_2 q_2 + \dots + F_n q_n$$

im ganzen Bereiche $\mathfrak{B}$ zu fragen. Der Bereich $\mathfrak{B}$ erstreckt sich ins Unendliche. Insbesondere ist $q_1 = q_2 = \dots = q_n = l$ ein Punkt aus $\mathfrak{B}$. Nun wird durch die Ungleichung

$$F_1 q_1 + F_2 q_2 + \dots + F_n q_n \leq l(F_1 + F_2 + \dots + F_n)$$

ein Teil von $\mathfrak{B}$ bestimmt, der ganz im Endlichen liegt und zum mindesten jenen speziellen Punkt enthält. In diesem Teile hat der Ausdruck (97) sicher einen bestimmten kleinsten Wert, welcher nun zugleich das Minimum dieses Ausdrucks (97) im ganzen Bereiche $\mathfrak{B}$ sein wird.

Dieser kleinste Wert von (97) in $\mathfrak{B}$ sei $3V^{\frac{2}{3}}$, und es sei $r_1, r_2, \dots, r_n$ ein Punkt in $\mathfrak{B}$, für den dieser Wert eintritt. Der Punkt (r_i) liegt dann jedenfalls auf der Begrenzung von $\mathfrak{B}$, es stellt also $\mathfrak{R}(r_i)$ ein Polyeder vom Volumen 1 vor. Nehmen wir mit diesem Polyeder diejenige Translation vor, wodurch sein Schwerpunkt in den Nullpunkt fällt, so treten an Stelle der Parameter r_i gewisse Werte $r_i + a\alpha_i + b\beta_i + c\gamma_i$; für diese hat dann aber der Ausdruck (97) genau denselben Wert, da wir für die Größen F_i die Gleichungen (93) als bestehend angenommen haben. Danach können wir auch von vornherein uns die Parameter r_i so beschaffen denken, daß der Schwerpunkt von $\mathfrak{R}(r_i)$ sich im Nullpunkte befindet. Alsdann sind notwendig die Größen $r_1, r_2, \dots, r_n$ sämtlich > 0 . Für alle Punkte (q_i) in $\mathfrak{B}$ gilt nun die Ungleichung

$$(98) \quad F_1 q_1 + F_2 q_2 + \dots + F_n q_n \geq 3V^{\frac{2}{3}}$$

und haben wir hierin also die Bedingung einer „Stützebene“ an $\mathfrak{B}$ durch den Punkt (r_i) , d. h. einer solchen Ebene, welche auf einer Seite von sich gar keinen Punkt von $\mathfrak{B}$ liegen hat.

Für das Polyeder $\mathfrak{R}(r_1, r_2, \dots, r_n)$ bedeute allgemein Φ_i den Flächeninhalt der Seitenfläche mit der äußeren Normale $(\alpha_i, \beta_i, \gamma_i)$, wobei wir

$\Phi_i = 0$ zu setzen hätten, falls eine solche Seitenfläche im Polyeder nicht wirklich vorkäme. Alsdann können wir zeigen, daß die Ebene

$$\frac{1}{3}(\Phi_1 q_1 + \Phi_2 q_2 + \cdots + \Phi_n q_n) = 1,$$

die jedenfalls durch den Punkt (r_i) geht, die *einzig*e Stützebene an $\mathfrak{B}$ durch diesen Punkt vorstellt, daß mithin die n Gleichungen

$$(99) \quad \Phi_i = \frac{F_i}{V^{\frac{2}{3}}} \quad (i = 1, 2, \dots, n)$$

gelten müssen. Denn hätten nicht diese n Gleichungen sämtlich statt, so könnten wir in der durch (98) dargestellten Stützebene an $\mathfrak{B}$ einen Punkt $(r_i + s_i)$ finden, für den

$$\frac{1}{3}(\Phi_1 s_1 + \Phi_2 s_2 + \cdots + \Phi_n s_n) > 0$$

ausfällt und noch alle Werte $r_i + s_i > 0$ sind. Nun ist $V(r_i) = 1$ und nach (91) das gemischte Volumen

$$V(\mathfrak{R}(r_i + s_i), \mathfrak{R}(r_i), \mathfrak{R}(r_i)) = 1 + \frac{1}{3} \sum \Phi_i s_i > 1;$$

infolgedessen wird das Volumen von $(1-t)\mathfrak{R}(r_i) + t\mathfrak{R}(r_i + s_i)$ dem Ausdrucke (60) zufolge bei hinreichend kleinen positiven Werten t sicher > 1 und umsomehr nach den oben gemachten Bemerkungen dann

$$V((1-t)r_i + t(r_i + s_i)) = V(r_i + ts_i) > 1.$$

Dieses könnte aber nicht der Fall sein, weil der Punkt $(r_i + ts_i)$ in einer Stützebene an $\mathfrak{B}$ liegt, somit gewiß nicht ein innerer Punkt von $\mathfrak{B}$ sein kann.

Damit ist in der Tat bewiesen, daß die n Gleichungen (99) gelten müssen. Setzen wir $V^{\frac{1}{3}} r_i = p_i$, so besitzt alsdann das Polyeder $\mathfrak{R}(p_i)$ zu den Normalen $(\alpha_i, \beta_i, \gamma_i)$ die Inhalte F_i der Seitenflächen, entspricht mithin genau den Forderungen unseres Satzes in 45.

§ 10. Bestimmung eines konvexen Körpers zu einer gegebenen Krümmungsfunktion.

49. Auf den soeben gewonnenen Satz über Polyeder gründen wir nun den Beweis des Satzes in 43. über die Existenz eines stetig gekrümmten konvexen Körpers zu einer gegebenen Krümmungsfunktion. Dabei wird uns namentlich die in (76) abgeleitete untere Grenze für die Differenz $V_1 - \sqrt[3]{V_0^2 V_3}$ von Nutzen sein.

Zunächst haben wir eine Bemerkung über gewisse Einteilungen auf der Kugelfläche $\mathfrak{E} (x^2 + y^2 + z^2 = 1)$ vorausszuschicken. Unter der *Winkeldistanz* zweier Punkte r und r^* auf $\mathfrak{E}$ wollen wir den Winkel $\angle ror^*$ der

Radien vom Nullpunkte o nach diesen Punkten verstehen. Es sei θ ein beliebiger positiver Wert, den wir $< \frac{\pi}{4}$ annehmen. Wir bestimmen auf $\mathfrak{E}$ sukzessive Punkte $r_1, r_2, \dots$ derart, daß die Winkeldistanz jedes folgenden Punktes von allen vorhergehenden $> \theta$ ist. Da die Oberfläche von $\mathfrak{E}$ eine endliche Größe hat, können wir eine solche Reihe von Punkten nicht unbegrenzt bilden, sondern *wir gelangen schließlich zu einer endlichen Reihe $r_1, r_2, \dots, r_n$ derart, daß nun für jeden beliebigen Punkt r auf $\mathfrak{E}$ unter den n Winkeldistanzen von r nach $r_1, r_2, \dots, r_n$ immer wenigstens eine $\leq \theta$ ausfällt.*

Es seien ξ_i, η_i, ζ_i die Koordinaten von r_i ($i = 1, 2, \dots, n$). Die n Ungleichungen

$$\xi_i x + \eta_i y + \zeta_i z \leq 1 \quad (i = 1, 2, \dots, n)$$

bestimmen alsdann ein Polyeder mit n Seitenflächen $\mathfrak{R}_i$, das ganz in der Kugel mit o als Mittelpunkt vom Radius $\frac{1}{\operatorname{tg} \theta}$ enthalten ist. Die Pyramide $o\mathfrak{R}_i$ mit o als Spitze und $\mathfrak{R}_i$ als Basis schneidet aus der Kugelfläche $\mathfrak{E}$ eine Partie $\mathfrak{E}_i$ heraus, den Bereich aller der Punkte r auf $\mathfrak{E}$, für welche die Winkeldistanz von dem betreffenden Punkte r_i nicht größer als von irgendeinem der anderen $n - 1$ Punkte r_j ($j \neq i$) ist. Es sei E_i der Flächeninhalt von $\mathfrak{E}_i$; das Gebiet $\mathfrak{E}_i$ enthält gewiß alle Punkte auf $\mathfrak{E}$ mit einer Winkeldistanz $\leq \frac{\theta}{2}$ von r_i und ist selbst ganz enthalten im Gebiet aller Punkte auf $\mathfrak{E}$ mit einer Winkeldistanz $\leq \theta$ von r_i ; daraus folgt

$$2\pi \left(1 - \cos \frac{\theta}{2}\right) < E_i < 2\pi (1 - \cos \theta).$$

Da überdies

$$E_1 + E_2 + \dots + E_n = 4\pi$$

ist, so haben wir

$$\frac{1}{2} n \left(1 - \cos \frac{\theta}{2}\right) < 1 < \frac{1}{2} n (1 - \cos \theta).$$

50. Es sei nun $F(\alpha, \beta, \gamma)$ die gegebene auf $\mathfrak{E}$ durchweg stetige und positive Funktion, welche die drei Gleichungen

$$(100) \quad \int \alpha F d\omega = 0, \quad \int \beta F d\omega = 0, \quad \int \gamma F d\omega = 0$$

erfüllt und als Krümmungsfunktion eines gewissen konvexen Körpers erkannt werden soll.

Verstehen wir unter $\alpha^*, \beta^*, \gamma^*$ ebenfalls einen beliebigen Punkt auf $\mathfrak{E}$, so stellt

$$(101) \quad S(\alpha^*, \beta^*, \gamma^*) = \frac{1}{2} \int |\alpha \alpha^* + \beta \beta^* + \gamma \gamma^*| F(\alpha, \beta, \gamma) d\omega$$

eine Funktion der Argumente $\alpha^*, \beta^*, \gamma^*$ dar, welche gleichfalls auf $\mathfrak{E}$ durchweg stetig und positiv sein wird, und besitzt diese Funktion dann auf $\mathfrak{E}$

ein bestimmtes *Minimum*, das wir S_0 nennen und das auch noch > 0 sein wird.

51. Wir denken uns jetzt eine Massenbelegung der Kugelfläche $\mathfrak{E}$ vorgenommen, wobei die Flächendichtigkeit an einem Punkte α, β, γ durch $F(\alpha, \beta, \gamma)$ dargestellt ist. Der Schwerpunkt der Belegung des Gebiets $\mathfrak{E}_i$ wird alsdann, da der Sektor $\circ\mathfrak{E}_i$ mit $\circ$ als Spitze und $\mathfrak{E}_i$ als Basis ein konvexer Körper ist und die Punkte auf $\mathfrak{E}_i$ eine Winkeldistanz $\leq \theta$ von r_i haben, ein Punkt q_i sein, der eine Entfernung $q_i > \cos \theta$ und < 1 von $\circ$ besitzt und so gelegen ist, daß der Strahl $\circ q_i$ in seiner Verlängerung einen Punkt p_i innerhalb $\mathfrak{E}_i$ trifft. Wir bezeichnen mit $\alpha_i, \beta_i, \gamma_i$ die Koordinaten von p_i , so daß $q_i\alpha_i, q_i\beta_i, q_i\gamma_i$ die von q_i sind. Setzen wir noch

$$(102) \quad \int_{(\mathfrak{E}_i)} F d\omega = \frac{F_i}{q_i},$$

so wird

$$(103) \quad \int_{(\mathfrak{E}_i)} \alpha F d\omega = \alpha_i F_i, \quad \int_{(\mathfrak{E}_i)} \beta F d\omega = \beta_i F_i, \quad \int_{(\mathfrak{E}_i)} \gamma F d\omega = \gamma_i F_i;$$

darin sind die vier Integrale über den Bereich $\mathfrak{E}_i$ zu erstrecken. Für die damit eingeführten n Richtungen $\alpha_i, \beta_i, \gamma_i$ und n positiven Größen F_i werden nunmehr auf Grund der Gleichungen (100) die Beziehungen

$$(104) \quad \sum \alpha_i F_i = 0, \quad \sum \beta_i F_i = 0, \quad \sum \gamma_i F_i = 0 \quad (i=1, 2, \dots, n)$$

statthaben.

Jeder Punkt auf $\mathfrak{E}$ besitzt von wenigstens einem der Punkte r_i eine Winkeldistanz $\leq \theta$, und sodann von dem zugehörigen Punkte p_i gewiß eine Winkeldistanz $< 2\theta$, weil immer r_i von p_i eine Winkeldistanz $< \theta$ hat. Da wir nun $2\theta < \frac{\pi}{2}$ vorausgesetzt haben, können hiernach die n Richtungen $\alpha_i, \beta_i, \gamma_i$ gewiß nicht sämtlich in einer Ebene liegen.

Nach dem Satze in 45. wird es nunmehr ein ganz bestimmtes Polyeder $\mathfrak{P}$ geben mit n Seitenflächen, den Richtungen $(\alpha_i, \beta_i, \gamma_i)$ als äußeren Normalen und den Größen F_i als Flächeninhalten dieser Seitenflächen, und zudem noch mit dem Schwerpunkte im Nullpunkte.

52. Wir haben jetzt noch einige allgemeinere Abschätzungen zur Sprache zu bringen.

Ist $\mathfrak{R}$ ein beliebiger konvexer Körper mit dem Schwerpunkte im Nullpunkte, H die Stützebenenfunktion von $\mathfrak{R}$, so wollen wir das Integral

$$(105) \quad \frac{1}{3} \int H(\alpha, \beta, \gamma) F(\alpha, \beta, \gamma) d\omega = J(\mathfrak{R})$$

schreiben. Es sei G das Maximum unter den Werten $H(\alpha, \beta, \gamma)$ auf $\mathfrak{E}$, also der Radius der kleinsten, den Körper $\mathfrak{R}$ ganz in sich enthaltenden Kugel mit dem Nullpunkte als Mittelpunkt, und etwa $G\alpha^*, G\beta^*, G\gamma^*$

ein solcher Punkt der Begrenzung von $\mathfrak{R}$, der vom Nullpunkte die Entfernung G hat. Nach der Definition der Stützebenenfunktion ist dann sicherlich immer

$$(106) \quad G(\alpha\alpha^* + \beta\beta^* + \gamma\gamma^*) \leq H(\alpha, \beta, \gamma).$$

Andererseits haben wir, weil der Schwerpunkt von $\mathfrak{R}$ sich im Nullpunkte befindet, mit Rücksicht auf die Ungleichung (59) stets

$$(107) \quad H(\alpha, \beta, \gamma) \leq 3H(-\alpha, -\beta, -\gamma).$$

Wir wenden nun auf das Integral (105) die Ungleichung (106) für alle diejenigen Richtungen α, β, γ an, wobei $\alpha\alpha^* + \beta\beta^* + \gamma\gamma^* \geq 0$ ist, und machen für die anderen Richtungen von (107) Gebrauch; beachten wir noch, daß zufolge der Gleichungen (100) das Integral

$$\frac{1}{2} \int |\alpha\alpha^* + \beta\beta^* + \gamma\gamma^*| F(\alpha, \beta, \gamma) d\omega,$$

erstreckt über die halben Kugelflächen von $\mathfrak{E}$, wo $\alpha\alpha^* + \beta\beta^* + \gamma\gamma^* \geq 0$ bzw. ≤ 0 gilt, beide Male denselben Wert hat, mithin, nach der in 50. erklärten Bedeutung der Größe S_0 , jedesmal $\geq \frac{1}{2} S_0$ sein muß, so folgt

$$(108) \quad J(\mathfrak{R}) \geq \frac{1}{3} \left(S_0 + \frac{1}{3} S_0 \right) G = \frac{4}{9} S_0 G.$$

Setzen wir

$$(109) \quad \int F(\alpha, \beta, \gamma) d\omega = O_0,$$

so wird andererseits

$$(110) \quad \frac{1}{3} O_0 G \geq J(\mathfrak{R}).$$

Bei der Annahme $F(\alpha, \beta, \gamma) = 1$, wobei die Relationen (100) jedenfalls erfüllt sind, geht auf diese Weise insbesondere

$$(111) \quad \frac{4\pi}{3} G \geq \frac{1}{3} \int H d\omega \geq \frac{4\pi}{9} G$$

hervor.

Setzen wir $H(\alpha_i, \beta_i, \gamma_i) = q_i$, so ist nach (91) das gemischte Volumen

$$(112) \quad V(\mathfrak{R}, \mathfrak{P}, \mathfrak{P}) = \frac{1}{3} (F_1 q_1 + F_2 q_2 + \cdots + F_n q_n).$$

Jeder Punkt α, β, γ auf $\mathfrak{E}$ besitzt von wenigstens einem Punkte $\alpha_i, \beta_i, \gamma_i$ eine Winkeldistanz $< 2\theta$, also eine geradlinige Distanz $< 2 \sin \theta$ und gilt alsdann nach der Ungleichung (6) in § 1 immer:

$$|H(\alpha, \beta, \gamma) - H(\alpha_i, \beta_i, \gamma_i)| < 2 \sin \theta \cdot G.$$

Mit Hilfe dieser Beziehung und der Formeln (102) gewinnen wir aus (105) und (112) einerseits:

$$(113) \quad J(\mathfrak{R}) > V(\mathfrak{R}, \mathfrak{P}, \mathfrak{P}) - 2 \sin \theta \cdot O_0 G;$$

andererseits:

$$(114) \quad \frac{1}{\cos \theta} V(\mathfrak{R}, \mathfrak{P}, \mathfrak{P}) > J(\mathfrak{R}) - 2 \sin \theta \cdot O_0 G.$$

53. Die abgeleiteten Ungleichungen benutzen wir zunächst, um in betreff der Ausdehnung des in 51. konstruierten Polyeders $\mathfrak{P}$ gewisse Grenzen nachzuweisen. Es sei $P(u, v, w)$ die Stützebenenfunktion von $\mathfrak{P}$, N das Maximum unter den Werten $P(\alpha, \beta, \gamma)$, ferner V das Volumen von $\mathfrak{P}$. Aus (111) entnehmen wir

$$(115) \quad \frac{4\pi}{3} N \geq \frac{1}{3} \int P(\alpha, \beta, \gamma) d\omega \geq \frac{4\pi}{9} N.$$

Die Oberfläche von $\mathfrak{P}$ ist

$$(116) \quad F_1 + F_2 + \dots + F_n < O_0 \quad \text{und} \quad > \cos \theta \cdot O_0.$$

Wenden wir nun die allgemeinen Ungleichungen $V_1^3 \geq V_0^2 V_3$, $V_2^3 \geq V_0 V_3^2$ für zwei beliebige konvexe Körper auf das Polyeder $\mathfrak{P}$ und die Kugel $\mathfrak{G}$ an, so erhalten wir mit Rücksicht hierauf:

$$(117) \quad \frac{1}{27} O_0^3 > \frac{4\pi}{3} V^2, \quad \frac{4\pi}{3} N^3 > V.$$

Aus (113) gewinnen wir

$$(118) \quad J(\mathfrak{P}) > V - 2 \sin \theta \cdot O_0 N$$

und aus (114) und (108):

$$(119) \quad \frac{V}{\cos \theta} > J(\mathfrak{P}) - 2 \sin \theta \cdot O_0 N \geq \left(\frac{4}{9} S_0 - 2 \sin \theta \cdot O_0 \right) N.$$

Mit Hilfe der zweiten Ungleichung in (117) folgt hieraus

$$(120) \quad V^{\frac{2}{3}} > \left(\frac{3}{4\pi} \right)^{\frac{1}{3}} \cos \theta \left(\frac{4}{9} S_0 - 2 \sin \theta \cdot O_0 \right).$$

Wir nehmen nun einen Winkel θ_0 so klein an, daß jedenfalls

$$\sin \theta_0 < \frac{2}{9} \frac{S_0}{O_0}$$

ist, und gewinnen dann aus (120) eine von θ unabhängige positive GröÙe V_0 und hernach aus der zweiten Ungleichung in (117) eine von θ unabhängige GröÙe N_0 derart, daß immer

$$(121) \quad V \geq V_0, \quad N_0 \geq N$$

statthat, wenn $\theta \leq \theta_0$ ist, was wir von nun an voraussetzen.

Das Polyeder $\frac{1}{V^{\frac{1}{3}}} \mathfrak{P}$ hat das Volumen 1. Aus (119) entnehmen wir

$$(122) \quad J\left(\frac{\mathfrak{P}}{V^{\frac{1}{3}}}\right) = \frac{1}{V^{\frac{1}{3}}} J(\mathfrak{P}) < \left(\frac{4\pi}{3}\right)^{\frac{2}{3}} \frac{1}{\cos \theta_0} N_0^2 + 2 \sin \theta_0 \cdot \frac{O_0 N_0}{V_0^{\frac{1}{3}}};$$

die hier rechts stehende GröÙe setzen wir $= \frac{4}{9} S_0 M_0$, dann ist also

$$(123) \quad J\left(\frac{\mathfrak{P}}{V^{\frac{1}{3}}}\right) < \frac{4}{9} S_0 M_0.$$

54. Es sei jetzt $\mathfrak{L}$ ein beliebiger konvexer Körper vom Volumen 1 und mit dem Nullpunkte als Schwerpunkt, $L(u, v, w)$ die Stützebenen-

funktion von $\mathfrak{L}$. Wir fragen, unter welchen Umständen sich

$$(124) \quad J(\mathfrak{L}) \leq J\left(\frac{\mathfrak{P}}{V^{\frac{1}{3}}}\right)$$

herausstellen kann. Nach der Formel (108) und infolge der Ungleichung (123) wird hierzu jedenfalls nötig sein, daß $\mathfrak{L}$ ganz im Inneren der Kugel vom Radius M_0 mit dem Nullpunkt als Mittelpunkt enthalten ist, daß also stets $L(\alpha, \beta, \gamma) < M_0$ gilt. Nach (113) ist sodann

$$(125) \quad J(\mathfrak{L}) > V(\mathfrak{L}, \mathfrak{P}, \mathfrak{P}) - 2 \sin \theta \cdot O_0 M_0.$$

Jetzt ziehen wir die Resultate des § 6 heran. Bezeichnen wir mit D das Maximum unter den Quotienten

$$(126) \quad \frac{V^{\frac{1}{3}} L(\alpha, \beta, \gamma)}{P(\alpha, \beta, \gamma)},$$

so gilt nach (76) die Ungleichung

$$(127) \quad \frac{V(\mathfrak{L}, \mathfrak{P}, \mathfrak{P})}{V^{\frac{2}{3}}} - 1 \geq \kappa \frac{(D-1)^6}{D^6},$$

worin κ die numerische Konstante $\frac{1}{2^{10} \cdot 3^4 \cdot 7^{\frac{1}{3}}}$ bedeutet. Aus (125), (124),

(119) und (121) schließen wir nun

$$(128) \quad \left(\frac{1}{\cos \theta} - 1\right) + 2 \sin \theta \cdot O_0 \left(\frac{M_0}{V_0^{\frac{2}{3}}} + \frac{N_0}{V_0}\right) > \kappa \frac{(D-1)^6}{D^6}.$$

Aus dieser Ungleichung entnehmen wir für $D-1$ eine obere Grenze, die nach Null konvergiert, wenn θ nach Null abnimmt.

Andererseits haben wir, wenn d das Minimum der Funktion (126) bedeutet, nach (78) und der an diese Ungleichung angeschlossenen Bemerkung:

$$(129) \quad D^2 - 1 \geq \kappa \frac{(1-d)^6}{d},$$

und hieraus ergibt sich weiter für $1-d$ eine obere Grenze, die mit θ zugleich nach Null konvergiert. Wir werden somit auch eine Größe ε , die zugleich mit θ nach Null konvergiert, angeben können, so daß

$$1 + \varepsilon \geq D, \quad d \geq \frac{1}{1 + \varepsilon}$$

gilt, und wir haben damit das Resultat erlangt:

Soll

$$J(\mathfrak{L}) \leq J\left(\frac{\mathfrak{P}}{V^{\frac{1}{3}}}\right)$$

ausfallen, so muß jedenfalls $\mathfrak{L}$ ganz in $(1 + \varepsilon) \frac{\mathfrak{P}}{V^{\frac{1}{3}}}$ enthalten sein und selbst das Polyeder $\frac{1}{1 + \varepsilon} \frac{\mathfrak{P}}{V^{\frac{1}{3}}}$ in sich enthalten, wobei ε eine gewisse vom Winkel θ abhängende Größe bedeutet, die mit nach Null abnehmendem θ ebenfalls nach Null konvergiert.

55. Dieses Resultat zeigt uns sofort (s. 9.), daß, wenn wir den Winkel θ nach Null abnehmen lassen, das Polyeder $\frac{\mathfrak{P}}{V^{\frac{1}{3}}}$ nach einem bestimmten konvexen Körper $\mathfrak{Q}$ als Grenze konvergieren muß, welchem die Eigenschaft zukommen wird, daß für ihn unter allen konvexen Körpern vom Volumen 1 und dem Nullpunkt als Schwerpunkt das Integral

$$J(\mathfrak{Q}) = \frac{1}{3} \int L(\alpha, \beta, \gamma) F(\alpha, \beta, \gamma) d\omega,$$

unter L die Stützebenenfunktion von $\mathfrak{Q}$ verstanden, den kleinsten Wert hat. Bezeichnen wir das betreffende Minimum dieses Integrals mit J , so konvergiert gleichzeitig $V^{\frac{2}{3}}$ nach J (s. (118) und (119)) und das Polyeder $\mathfrak{P}$ nach dem Körper $\mathfrak{R} = J^{\frac{1}{3}} \mathfrak{Q}$.

Ist jetzt $\mathfrak{R}'$ ein beliebiger konvexer Körper, H' seine Stützebenenfunktion, G' das Maximum der Werte $H'(\alpha, \beta, \gamma)$, so folgt aus (113), (114) und (110):

$$|J(\mathfrak{R}') - V(\mathfrak{R}', \mathfrak{P}, \mathfrak{P})| < \left(\frac{1}{3} (1 - \cos \theta) + 2 \sin \theta \right) O_0 G'.$$

Da nun für ein nach Null abnehmendes θ die Größe $V(\mathfrak{R}', \mathfrak{P}, \mathfrak{P})$ nach $V(\mathfrak{R}', \mathfrak{R}, \mathfrak{R})$ konvergiert, so ersehen wir hieraus, daß allgemein die Darstellung

$$V(\mathfrak{R}', \mathfrak{R}, \mathfrak{R}) = J(\mathfrak{R}'), \quad \text{d. i.} \quad = \frac{1}{3} \int H' F d\omega$$

gilt. Danach ist in der Tat der gefundene konvexe Körper $\mathfrak{R}$ ein stetig gekrümmter mit $F(\alpha, \beta, \gamma)$ als Krümmungsfunktion, mithin der Beweis für den Satz in 43. vollständig erbracht.

56. Wir können diesen Satz über die Bestimmung eines konvexen Körpers zu einer gegebenen Krümmungsfunktion F auch auf Fälle ausdehnen, wo diese vorgelegte Funktion F nicht durchweg stetig ist. Insbesondere lassen sich alle konvexen Körper, welche in der Regel *konstante* positive Krümmung und nur an singulären Stellen unendliche Krümmung besitzen, durch folgende Aussage charakterisieren:

Es seien auf der Kugelfläche $\mathfrak{E}$ beliebige Partien $\mathfrak{N}$ abgegrenzt, denen ein bestimmter und von Null verschiedener Flächeninhalt zukommt und so, daß der Schwerpunkt dieser Partien $\mathfrak{N}$ für sich, wie der der ganzen Kugelfläche $\mathfrak{E}$, im Nullpunkte liegt. Alsdann gibt es stets einen und nur einen konvexen Körper $\mathfrak{R}$ mit dem Nullpunkte als Schwerpunkt, derart daß für jeden beliebigen konvexen Körper $\mathfrak{R}'$ die Darstellung

$$V(\mathfrak{R}', \mathfrak{R}, \mathfrak{R}) = \frac{1}{3} \int_{(\mathfrak{N})} H' d\omega$$

gilt, wo H' die Stützebenenfunktion von $\mathfrak{R}'$ bedeutet und das Integral nur über die Partien $\mathfrak{N}$ der Kugelfläche $\mathfrak{E}$ zu erstrecken ist.

XXVII.

Über die Körper konstanter Breite.

(In russischer Sprache erschienen in: Mathematische Sammlung (Matematičeskij Sbornik), Moskau, Band 25, S. 505—508.)

§ 1.

Unter der *Breite eines Körpers in einer gegebenen Richtung* soll der Abstand der zwei zu dieser Richtung normalen Stützebenen des Körpers voneinander verstanden werden. Als *Stützebene* des Körpers wird hier jede Ebene bezeichnet, welche an den Körper in wenigstens einem Punkte anstößt, ihn aber nicht durchsetzt.

Ein Körper, der in allen Richtungen gleiche Breite hat, soll von *konstanter Breite* heißen.

§ 2.

Es sei irgendein Körper K vorgelegt. Wir verwenden rechtwinklige Koordinaten x, y, z mit einem Punkte O im Inneren von K als Anfangspunkt, und führen zur Festlegung der Punkte α, β, γ auf der Kugelfläche vom Radius 1 um O Polarkoordinaten ϑ, ψ ein, so daß

$\alpha = \sin \vartheta \cos \psi, \quad \beta = \sin \vartheta \sin \psi, \quad \gamma = \cos \vartheta \quad (0 \leq \vartheta \leq \pi, \quad 0 \leq \psi \leq 2\pi)$ ist.

Der Abstand derjenigen Stützebene an K , welche α, β, γ als die vom Körper abgewandte Normalenrichtung zeigt, vom Nullpunkt heiße $H(\vartheta, \psi)$.

Der zu einer Richtung $\alpha, \beta, \gamma(\vartheta, \psi)$ entgegengesetzten Richtung $-\alpha, -\beta, -\gamma$ entsprechen dann die Werte $\pi - \vartheta, \psi + \pi \pmod{2\pi}$ dieser Polarkoordinaten.

Die Breite von K in der Richtung α, β, γ wird alsdann durch

$$B(\vartheta, \psi) = H(\vartheta, \psi) + H(\pi - \vartheta, \psi + \pi)$$

dargestellt. Dieser Ausdruck bleibt ungeändert, wenn wir α, β, γ durch die entgegengesetzte Richtung $-\alpha, -\beta, -\gamma$ ersetzen.

Denken wir uns $H(\vartheta, \psi)$ in eine Reihe nach Kugelflächenfunktionen entwickelt,

$$H(\vartheta, \psi) = Y_0 + Y_1(\vartheta, \psi) + Y_2(\vartheta, \psi) + \dots,$$

worin Y_0 eine Konstante und $Y_m(\vartheta, \psi)$ eine Kugelfunktion m^{ter} Ordnung vorstellt, so folgt

$$H(\pi - \vartheta, \psi + \pi) = Y_0 - Y_1(\vartheta, \psi) + Y_2(\vartheta, \psi) + \dots$$

und

$$B(\vartheta, \psi) = 2Y_0 + 2Y_2(\vartheta, \psi) + 2Y_4(\vartheta, \psi) + \dots$$

Damit K einen Körper konstanter Breite vorstellt, ist hiernach notwendig und hinreichend, daß die Terme gerader Ordnung $Y_2(\vartheta, \psi)$, $Y_4(\vartheta, \psi)$, $\dots$ in der Entwicklung von $H(\vartheta, \psi)$ sämtlich identisch verschwinden.

§ 3.

Unter dem *Umfange* eines Körpers in bezug auf eine gegebene Richtung wollen wir den *Umfang des Querschnittes* bei demjenigen Zylinder verstehen, der von allen der betreffenden Richtung parallelen Stützebenen des Körpers umhüllt wird.

Ein Körper, der in bezug auf alle Richtungen gleichen Umfang hat, soll von *konstantem Umfange* heißen.

Wir werden hier den Satz beweisen:

Die Körper konstanter Breite und die Körper konstanten Umfanges sind miteinander identisch.

§ 4.

Der Umfang des Körpers K speziell in bezug auf die Richtung der z -Achse ist der Umfang der in der xy -Ebene von den sämtlichen Geraden

$$x \cos \psi + y \sin \psi = H\left(\frac{\pi}{2}, \psi\right) = h(\psi)$$

umhüllten geschlossenen konvexen Kurve. Für die Punkte x, y dieser Kurve haben wir neben dieser ersten Gleichung noch die andere

$$-x \sin \psi + y \cos \psi = \frac{\partial h(\psi)}{\partial \psi},$$

so daß für sie

$$x = h \cos \psi - \frac{\partial h}{\partial \psi} \sin \psi, \quad y = h \sin \psi + \frac{\partial h}{\partial \psi} \cos \psi,$$

$$dx = -\left(h + \frac{\partial^2 h}{\partial \psi^2}\right) \sin \psi d\psi, \quad dy = \left(h + \frac{\partial^2 h}{\partial \psi^2}\right) \cos \psi d\psi$$

gilt und also der Umfang jener Kurve durch

$$\int_0^{2\pi} \left(h + \frac{\partial^2 h}{\partial \psi^2}\right) d\psi = \int_0^{2\pi} h d\psi$$

dargestellt wird.

Setzen wir nun

$$Y_m(\vartheta, \psi) = A_0 P^{(m)}(\cos \vartheta) + \sum_{\nu=1}^m (A_\nu \cos \nu \psi + B_\nu \sin \nu \psi) P_\nu^{(m)}(\cos \vartheta)$$

an, wobei

$$P^{(m)}(1) = 1, \quad P^{(2l-1)}(0) = 0, \quad P^{(2l)}(0) = (-1)^l \frac{1 \cdot 3 \cdot 5 \cdots (2l-1)}{2 \cdot 4 \cdot 6 \cdots 2l}$$

ist,

$$P_\nu^{(m)}(1) = 0 \quad (\nu = 1, \dots, m),$$

so erkennen wir, daß

$$\int_0^{2\pi} Y_m\left(\frac{\pi}{2}, \psi\right) d\psi$$

bei ungeradem m verschwindet und bei geradem m sich als das Produkt aus dem Werte der Funktion $Y_m(\vartheta, \psi)$ für den Pol $\vartheta = 0$ in eine gewisse von m abhängige Konstante $\tilde{\omega}_m = 2\pi P^{(m)}(0)$ ergibt.

Indem wir dieses in bezug auf die Richtung der z -Achse gewonnene Resultat auf eine beliebige Richtung übertragen, sehen wir, daß aus der angenommenen Entwicklung für $H(\vartheta, \psi)$ sich für den Umfang von K in bezug auf eine beliebige Richtung (ϑ, ψ) die folgende Entwicklung

$$U(\vartheta, \psi) = 2\pi Y_0 + \tilde{\omega}_2 Y_2(\vartheta, \psi) + \tilde{\omega}_4 Y_4(\vartheta, \psi) + \dots$$

herausstellt.

Damit K ein Körper konstanten Umfanges ist, finden wir hiernach notwendig und hinreichend, daß in der Entwicklung von $H(\vartheta, \psi)$ nach Kugelfunktionen die Terme $Y_2, Y_4, \dots$ sämtlich Null sind.

Ein Vergleich dieses Resultats mit dem vorhin gewonnenen über die Körper konstanter Breite zeigt in der Tat, daß ein Körper konstanter Breite jedesmal zugleich ein Körper konstanten Umfanges ist und umgekehrt.

ZUR PHYSIK

XXVIII.

Über die Bewegung eines festen Körpers in einer Flüssigkeit.

(Sitzungsberichte der K. Preußischen Akademie der Wissenschaften zu Berlin.
Band XL. 1888. S. 1095—1110.)

(Vorgelegt von Herrn von Helmholtz.)

Gustav Kirchhoff*) hat aus dem Hamiltonschen Prinzip Differentialgleichungen für die Bewegung eines festen Körpers in einer Flüssigkeit hergeleitet, und dieselben Gleichungen hat Sir William Thomson**) mit Hilfe seiner allgemeinen Bestimmung der nach Impulsen eintretenden Zustände gefunden. Es liegen diesen Gleichungen die Vorstellungen zugrunde, daß der Körper von unveränderlichem Gefüge, die Flüssigkeit homogen, inkompressibel und reibungslos ist, ferner, daß die Flüssigkeit nach allen Richtungen sich in die Unendlichkeit erstreckt, dort überall ruht, und daß an jedem ihrer Punkte ein einwertiges Geschwindigkeitspotential existiert. Diese Vorstellungen halte ich im folgenden fest, und ich gebe für den Fall, daß keine Kräfte wirken, eine Reduktion jener Gleichungen auf ein Problem, das mit dem der kürzesten Linien auf einem Ellipsoide Ähnlichkeit hat. Bisher sind, selbst für diesen elementarsten Fall, nur einzelne Integrale, partikuläre Lösungen, und Folgerungen, die sich auf Körper von spezieller Gestalt und Massenverteilung beziehen, bekannt; hier werde ich keinerlei Beschränkungen hinsichtlich des Körpers eintreten lassen.

Zunächst suche ich in § 1 diejenige Zerlegung der Bewegung eines Körpers in einer Flüssigkeit auf, welche das Analogon ist zu der Zerlegung der Bewegung eines Körpers im leeren Raume in die Bewegung des Trägheitsmittelpunktes des Körpers und die Rotation um diesen Punkt. In § 2 ist die Bedeutung der Differentialgleichungen von Kirchhoff und Thomson auseinandergesetzt, und werden diejenigen Bewegungszustände

*) Crelles Journal, Bd. 71, S. 237—262. — Auch in den Vorlesungen über Mechanik, XIX. Vorl.

**) Philosophical Magazine, 1871, Vol. 42, pp. 362—366.

eines Körpers in einer Flüssigkeit bestimmt, die sich, ohne Einfluß von Kräften, unverändert erhalten. Verschiedene von den Resultaten dieses Paragraphen finden sich bereits in den Aufsätzen der Herren H. Lamb*) und Th. Craig**); aber die Entwicklung derselben geschieht hier einfacher und übersichtlicher. In § 3 leite ich die Bewegung des Körpers — so wie sie vom Körper gesehen sich darstellt — für den Fall, daß keine Kräfte wirken, aus dem einen, der Bewegung eines Trägheitsmittelpunktes analogen Teile allein ab. Indem ich dann auch aus dem Ausdrucke des Hamiltonschen Prinzips die Koordinaten des anderen Teiles herausschaffe, tritt eine neue und sehr anschauliche Minimaleigenschaft zutage. Diese ermöglicht schließlich eine Anwendung der Sätze von C. G. J. Jacobi***) über diejenigen Probleme der Mechanik, welche nur zwei zu bestimmende Größen enthalten.

Eine ganz ähnliche und nicht minder bemerkenswerte Gelegenheit, von diesen wichtigen Sätzen Nutzen zu ziehen, bietet sich, worauf ich indessen hier nicht eingehe, in dem Probleme der Rotation eines, der *Schwere* unterworfenen Körpers um einen festen Punkt; ob die Rotation im leeren Raume oder in einer unendlichen, schweren Flüssigkeit vor sich geht, macht für die mathematische Behandlung nur insoweit einen Unterschied aus, als im letzteren Falle von den drei, auf den festen Punkt bezüglichen Hauptträgheitsmomenten auch eines größer als die Summe der beiden anderen sein kann, was im ersteren Falle ausgeschlossen ist.

§ 1.

In fester Verbindung mit dem betrachteten Körper sei ein rechtwinkliges Koordinatensystem angenommen; anstatt von der Bewegung des Körpers können wir von der Bewegung dieses Systems sprechen. Es seien u, v, w die augenblicklichen Geschwindigkeiten des Anfangspunktes dieses Systems parallel zu den Achsen desselben, p, q, r die augenblicklichen Drehungsgeschwindigkeiten des Systems um diese Achsen. Über den Sinn positiver Drehungen sei die gewöhnliche Festsetzung getroffen, derzufolge die Ausdrücke für die Geschwindigkeiten eines Punktes, dessen Koordinaten x, y, z sind, lauten:

$$(1) \quad u + zq - yr, \quad v + xr - zp, \quad w + yp - xq.$$

Mit unseren Voraussetzungen über die Ausdehnung und die Bewegung der Flüssigkeit erreichen wir, daß durch die Werte von u, v, w, p, q, r

*) Proceedings of the London Mathematical Society, 1877, VIII. pp. 273—286.

**) American Journal of Mathematics, 1879. II. pp. 162—171.

***) Vorlesungen über Dynamik, XXII. Vorl. — Ges. Werke, Supplementband, 1884, S. 175.

jedesmal auch der Bewegungszustand der Flüssigkeit völlig bestimmt ist, und daß wir uns den Zustand des ganzen, aus Körper und Flüssigkeit bestehenden Systems von der Ruhe aus momentan, durch einen gewissen, auf den Körper ausgeübten Impuls, entstanden denken können. Infolge dieser Umstände ist dann die gesamte lebendige Kraft des Körpers und der Flüssigkeit, die wir mit T bezeichnen, eine homogene Funktion zweiten Grades von u, v, w, p, q, r mit konstanten Koeffizienten, deren Werte von der Gestalt des Körpers, der Masse dieses und ihrer Verteilung, sowie von der Dichtigkeit der Flüssigkeit abhängig sind; und der in Frage kommende Impuls, den wir kurz den augenblicklichen *Impuls der Bewegung* nennen, und der eine, T gleiche Arbeit zu leisten haben würde, ist äquivalent einer im Anfangspunkte der Koordinaten angreifenden impulsiven Kraft, deren Komponenten $\frac{\partial T}{\partial u}, \frac{\partial T}{\partial v}, \frac{\partial T}{\partial w}$ sind, verbunden mit einem

impulsiven Kräftepaare von den Komponenten $\frac{\partial T}{\partial p}, \frac{\partial T}{\partial q}, \frac{\partial T}{\partial r}$. Diese Differentialquotienten sind hier nur als Symbole für gewisse lineare Ausdrücke in u, v, w, p, q, r aufzufassen, die ich auch der Reihe nach mit u, v, w, p, q, r bezeichne.

Zerlegung der lebendigen Kraft. In derselben Beziehung, wie die letzten sechs Größen zum Anfangspunkte der Koordinaten, stehen, vermöge der Ausdrücke (1), zu einem Punkte, dessen Koordinaten x, y, z sind, die Größen:

$$u, v, w, p + zv - yw, q + xw - zu, r + yu - xv.$$

Da hiernach die Bedeutung der impulsiven Einzelkräfte u, v, w — ähnlich wie die der Drehungsgeschwindigkeiten p, q, r — von der Wahl des Anfangspunktes völlig unabhängig ist und nur von den Richtungen der Achsen abhängt, so führe ich diese Größen u, v, w als neue Variable an Stelle von u, v, w ein. Das kann geschehen, da die in Frage kommende Determinante niemals verschwindet, weil T eine wesentlich positive Form ist. Die Form T mit sechs Variablen zerfällt aber dadurch in eine Form der drei Variablen u, v, w , und eine Form der drei Variablen p, q, r , die ihrerseits beide wesentlich positiv sein müssen; ich schreibe:

$$T = E(u, v, w) + G(p, q, r).$$

Die Größen u, v, w werden lineare Funktionen in u, v, w, p, q, r mit konstanten Koeffizienten. Aus den Ausdrücken (1), welche gewisse Funktionen eines beliebigen Punktes vorstellen, geht hervor, daß derjenige Punkt im Körper, dessen Koordinaten

$$\frac{1}{2} \left(\frac{\partial w}{\partial q} - \frac{\partial v}{\partial r} \right), \quad \frac{1}{2} \left(\frac{\partial u}{\partial r} - \frac{\partial w}{\partial p} \right), \quad \frac{1}{2} \left(\frac{\partial v}{\partial p} - \frac{\partial u}{\partial q} \right)$$

sind, eine von der Wahl des Koordinatensystems völlig unabhängige Bedeutung hat.

In diesen Punkt, den ich im Hinblick auf eine später auseinanderzusetzende Eigenschaft das *Zentrum der Hauptachsen des Körpers* nenne, soll für die Folge stets der Anfangspunkt des im Körper festen Koordinatensystems gelegt sein, d. h. ich setze die vorstehenden Differenzen gleich Null voraus. Dann lassen sich die Differenzen

$$u - \frac{\partial E}{\partial u}, \quad v - \frac{\partial E}{\partial v}, \quad w - \frac{\partial E}{\partial w},$$

die nicht mehr von u, v, w abhängen, als partielle Differentialquotienten nach p, q, r von einer gewissen quadratischen Form in p, q, r darstellen, die ich mit $-F(p, q, r)$ bezeichne; die Differenzen

$$p - \frac{\partial G}{\partial p}, \quad q - \frac{\partial G}{\partial q}, \quad r - \frac{\partial G}{\partial r}$$

sind hernach gleich den partiellen Differentialquotienten nach u, v, w von $+F(u, v, w)$.

Ein identisches Verschwinden dieser Form F — wie es beispielsweise immer eintritt, wenn der Körper der Gestalt und Verteilung der Masse nach ein Rotationskörper ist, oder wenn die Dichtigkeit der Flüssigkeit Null ist, d. h. die Bewegung im leeren Raume erfolgt — zeigt an, daß der betrachtete Körper hinsichtlich des Ausdrucks der lebendigen Kraft T den Charakter solcher Körper trägt, die der Gestalt und Massenverteilung nach einen *Mittelpunkt* aufweisen.

Der Ort des Zentrums der Hauptachsen und die drei Formen E, F, G mit je drei Variablen involvieren zusammen dieselbe Zahl von Konstanten, nämlich 21, wie die eine Form T mit sechs Variablen.

Ebenso wie die gesamte lebendige Kraft, erscheint die Bewegung des Körpers in zwei, von der Wahl eines Koordinatensystems völlig unabhängige Teile zerlegt, nämlich in eine *Verschiebungsgeschwindigkeit* ($p = 0, q = 0, r = 0$) und eine Bewegung, deren Impuls ein *impulsives Kräftepaar* ist ($u = 0, v = 0, w = 0$). Im leeren Raume würde der erste Teil die Bewegung des Trägheitsmittelpunktes, der zweite die Rotation um diesen Punkt vorstellen. Die Verschiebungsgeschwindigkeit hat zu Komponenten: $\frac{\partial E}{\partial u}, \frac{\partial E}{\partial v}, \frac{\partial E}{\partial w}$, das impulsive Kräftepaar: $\frac{\partial G}{\partial p}, \frac{\partial G}{\partial q}, \frac{\partial G}{\partial r}$.

Um diese Zerlegung weiter zu verfolgen, denke ich mir für einen Moment die Koordinatenachsen in Hauptachsen einer Fläche zweiten Grades $F(x, y, z) = \text{konst.}$ gelegt, so daß für $F(x, y, z)$ ein Ausdruck

$$\frac{1}{2} (ax^2 + by^2 + cz^2)$$

bestehe. Dann soll eine Linie, die in einer Richtung, deren Kosinus ξ, η, ζ sind, von dem Punkte

(2) $x = (b - c)\xi\eta$, $y = (c - a)\xi\xi$, $z = (a - b)\xi\eta$ ($\xi^2 + \eta^2 + \xi^2 = 1$)
 oder von dem, diesem Punkte in bezug auf das Zentrum der Hauptachsen gegenüberliegenden Punkte $-x$, $-y$, $-z$ ausgeht und sich ins Unendliche erstrecken mag, der zu der Richtung $\xi:\eta:\xi$ gehörende *erste, bzw. zweite Radius des Körpers* heißen. Als gleichwertig mit $\xi:\eta:\xi$ können hier offenbar nur solche Proportionen $m\xi:m\eta:m\xi$ gelten, in welchen m ein positiver Faktor ist. Zwei zu entgegengesetzten Richtungen gehörende erste, bzw. zweite Radien greifen stets an demselben Punkte an; die Linie, die sie gemeinsam bestimmen, soll ein *erster, bzw. zweiter Durchmesser des Körpers* heißen.

Man bestätigt leicht mit Hilfe der Ausdrücke (1), daß durch das impulsive Kräftepaar $\frac{\partial G}{\partial p}, \frac{\partial G}{\partial q}, \frac{\partial G}{\partial r}$ eine Schraubenbewegung um den zu der Richtung $p:q:r$ gehörenden ersten Radius entsteht, wobei die Geschwindigkeit der Drehung die Größe $+\sqrt{p^2 + q^2 + r^2}$ und die der Steigung die Größe $\frac{-2F(p, q, r)}{p^2 + q^2 + r^2}$ erlangt. Ganz analog reduziert sich der Impuls der Verschiebungsgeschwindigkeit $\frac{\partial E}{\partial u}, \frac{\partial E}{\partial v}, \frac{\partial E}{\partial w}$ auf eine impulsive Kraft $+\sqrt{u^2 + v^2 + w^2}$ längs dem zu der Richtung $u:v:w$ gehörenden zweiten Radius, und ein impulsives Kräftepaar $\frac{2F(u, v, w)}{u^2 + v^2 + w^2} \sqrt{u^2 + v^2 + w^2}$ um diesen Radius.

Die zu den sämtlichen Richtungen gehörenden ersten (oder zweiten) Durchmesser bestimmen im Körper ein, noch von der Dichtigkeit der Flüssigkeit abhängendes, Strahlensystem zweiter Klasse mit ausgezeichneten metrischen Eigenschaften: Die Mittelpunkte der Strahlen, identisch mit den Punkten (2), sind zugleich die Fußpunkte ihrer kürzesten Entfernungen vom Zentrum der Hauptachsen; sie bilden in ihrer Gesamtheit eine geschlossene Steinersche Fläche, welche sich auf die Fläche eines Kreises oder gar auf das Zentrum der Hauptachsen zusammenzieht, wenn $F = \text{konst.}$ eine Rotationsfläche oder eine Kugel wird. Sieht man ferner irgend drei zueinander senkrechte erste (oder zweite) Durchmesser als nichtzusammenhängende Kanten eines Parallelepipeds an, so fällt der Mittelpunkt dieses in das Zentrum der Hauptachsen.

§ 2.

Man kann zunächst aus dem Hamiltonschen Prinzipie sehr einfach die am Anfange von § 1 auseinandergesetzte Bedeutung der Differentialquotienten der lebendigen Kraft T nach den Geschwindigkeiten erschließen; dann sind die Differentialgleichungen von Kirchhoff und Thomson nur ein Ausdruck für die fast selbstverständliche Tatsache, daß der augen-

blickliche Impuls der Bewegung in jedem Zeitelement dt genau um den, von den gerade vorhandenen Kräften während dt ausgeübten Impuls zunimmt, vorausgesetzt, daß ein jeder Impuls nicht allein der Größe nach, sondern auch nach Richtung und Lage im Raume geschätzt wird.

Wirken keine Kräfte, und auch nur in solchem Falle wird daher der Impuls der Bewegung einen unveränderlichen Ausdruck im Raume haben. Diese Bedingung bestimmt den ganzen Bewegungsvorgang; denn man erkennt aus ihr, welche Änderungen der Geschwindigkeiten mit jeder Ortsänderung des Körpers einhergehen müssen.

Der augenblickliche Impuls der Bewegung besitzt, wie jeder Impuls, eine einzige, völlig sichere Darstellung, sei es als nichtverschwindende impulsive Kraft mit einem, zu der Kraft *senkrechten* impulsiven Kräftepaare, sei es bloß als impulsives Kräftepaar. Unter der *Achse des Impulses* versteht man in dem einen Falle die der Richtung und Lage nach völlig bestimmte Linie, in welcher die impulsive Einzelkraft wirkt, in dem anderen Falle die nur der Richtung nach bestimmte Achse des alsdann allein vorhandenen impulsiven Kräftepaars. Die *Größe des Impulses* sei definiert durch den Inbegriff zweier Größen J und J_1 , von denen die erste, J , die absolute Größe der impulsiven Einzelkraft, und die andere, J_1 , die im Sinne der Achsenrichtung des Impulses gemessene Größe des Moments des impulsiven Kräftepaars bezeichnen soll.

Wenn keine Kräfte wirken, so muß also die Größe des Impulses konstant und seine Achse im Raume unveränderlich sein. Daß alsdann auch die gesamte lebendige Kraft T konstant sein wird, leuchtet aus dem Satze von der Erhaltung der lebendigen Kraft ein.

Den analytischen Ausdruck dieser Bedingungen findet man in folgender Weise, wobei man den Fall $J = 0$ als Grenzfall eines unendlich kleinen J auffassen kann. Zunächst ist für die Größe des Impulses:

$$(3) \quad u^2 + v^2 + w^2 = J^2, \quad J \geq 0$$

$$(4) \quad up + vq + wr = JJ_1,$$

(wenn $J = 0$, d. h. $u = 0, v = 0, w = 0$ ist, so hat man: $p^2 + q^2 + r^2 = J_1^2, J_1 \geq 0$). Die Achse des Impulses besitzt in bezug auf das im Körper angenommene Koordinatensystem die Gleichungen:

$$(5) \quad p + zv - yw : q + xw - zu : r + yu - xv : J_1 = u : v : w : J$$

(oder ist im Falle $J = 0$ durch $p : q : r$ der Richtung nach bestimmt); nach Verlauf eines Zeitelements dt haben sich hierin u, v, w, p, q, r um ihre Differentiale während dt geändert; da aber die Achse des Impulses dieselbe geblieben sein soll, und nur der Ort des Koordinatensystems sich geändert hat, so müssen die Gleichungen dieser Linie nach Verlauf von dt auch zum Vorschein kommen (bis auf unendlich kleine Größen zweiter Ordnung),

wenn man in (5) nur zu x, y, z die in dt multiplizierten Geschwindigkeiten (1) hinzufügt. Hält man die Bedingung, daß die auf diese zweierlei Arten aus (5) entstehenden Gleichungen dieselben Punkte x, y, z definieren sollen, mit der anderen zusammen, daß J und J_1 Konstanten, d. h. die Differentialquotienten der linken Seiten von (3) und (4) nach t Null sein sollen, so hat man Beziehungen genug, um den Differentialquotienten jeder der Größen u, v, w, p, q, r nach t als Funktion eben dieser Größen darzustellen, worauf die Differentialgleichungen von Kirchhoff und Thomson hinauslaufen.

Stationäre Bewegungen. Wie der Körper auch beschaffen sein mag, so gibt es stets für ihn *stationäre* Bewegungszustände, d. h. es gibt Wertsysteme der Geschwindigkeiten u, v, w, p, q, r , die, einmal erzeugt, unverändert bestehen, so lange als keine Kräfte wirken. Eine Bewegung, die mit Geschwindigkeiten solcher Art beginnt, setzt sich als gleichförmige Schraubenbewegung fort.

Die charakteristische Bedingung für stationäre Bewegungszustände ist die, daß die Achse der Bewegung (nämlich der durch die Geschwindigkeiten bestimmten Schraubenbewegung) und die Achse des Impulses der Bewegung der Lage nach zusammenfallen müssen.

Denn soll irgendein Bewegungszustand eine Weile andauern, der Körper also eine gleichförmige Schraubenbewegung ausführen, so muß der Impuls der Bewegung diese Schraubenbewegung sozusagen mitmachen; dabei bleibt natürlich seine Größe ungeändert, seine Achse aber ist nur dann fest, wenn sie in die Achse der Bewegung fällt.

Die genannte Bedingung ist, falls eine der Achsen nur der Richtung nach definiert, also keine Drehungsgeschwindigkeit oder keine impulsive Einzelkraft vorhanden ist, dahin aufzufassen, daß den Achsen gleiche oder entgegengesetzte Richtungen zukommen müssen.

Für die Achse des Impulses ist bereits ein analytischer Ausdruck angegeben; die Achse der Bewegung hat, wenn die Drehungsgeschwindigkeit nicht Null ist, die Gleichungen:

$$u + zq - yr : v + xr - zp : w + yp - xq = p : q : r,$$

anderenfalls ist sie durch $u : v : w$ der Richtung nach definiert. Damit diesen zweierlei Bestimmungen *eine* Linie genügen könne, ist notwendig und hinreichend, daß mit irgendwelchen Werten von λ und μ die Beziehungen bestehen:

$$u : v : w : 1 = p : q : r : \lambda = u - \lambda p : v - \lambda q : w - \lambda r : \mu,$$

d. i.

$$0 = \delta \left(T - \lambda J J_1 - \frac{\mu}{2} J^2 \right) \left[= \delta \left(T + \frac{\lambda^2}{2\mu} J_1^2 \right), \text{ wenn } J = 0 \text{ ist} \right].$$

Danach haben in stationären Bewegungen die Geschwindigkeiten solche

Werte, daß die, bei festgehaltener Größe des Impulses genommene erste Variation der lebendigen Kraft identisch verschwindet (dasselbe ist von der ersten Variation bei festgehaltener Größe der Schraubenbewegung auszusagen).

Um einen Überblick über sämtliche überhaupt existierenden stationären Bewegungen eines Körpers zu erhalten, kann man sie nach den Werten ihres λ anordnen; λ bedeutet für sie das Verhältnis ihrer Drehungsgeschwindigkeit zu ihrer impulsiven Einzelkraft, und zwar, wenn dasselbe nicht Null oder unendlich ist, mit positivem oder negativem Vorzeichen, je nachdem ihre zwei Achsen der Bewegung und des Impulses gleiche oder entgegengesetzte Richtung haben:

Die Verhältnisse der Geschwindigkeiten in den stationären Bewegungen, welche zu einem festen Werte von λ gehören, erfüllen die Proportion $u:v:w:1 = p:q:r:\lambda$ und die Bedingung, daß die durch $u:v:w$ (oder $p:q:r$) bestimmte Richtung die einer Hauptachse der Fläche zweiten Grades

$$(6) \quad T - \lambda J J_1 = E(u, v, w) - 2\lambda F(u, v, w) - \lambda^2 G(u, v, w) = \text{konst.}$$

sein muß, wenn man u, v, w als rechtwinklige Koordinaten eines Punktes (an Stelle von x, y, z) ansieht. Der Parameter λ kann hier jeden reellen Wert haben, und für jeden Wert von λ gehören so zu jeder Hauptachse der Fläche (6) je zwei, einander entgegengesetzte stationäre Bewegungen mit beliebigem Werte der lebendigen Kraft T .

Die Gesamtheit der Hauptachsen aller Flächen (6) bestimmt im Körper für gewöhnlich einen Kegel sechster Ordnung. Auf der, diesem Kegel parallelen, von den *Doppelachsen* sämtlicher stationären Bewegungen gebildeten geradlinigen Fläche haben irgend drei, zu demselben λ gehörende und zueinander senkrechte Doppelachsen stets eine solche Lage, daß der Mittelpunkt des Parallelepipedums, für welches sie nichtzusammenhängende Kanten sind, in unseren Koordinatenanfangspunkt fällt.

$\lambda = \infty$. Der letzten Eigenschaft ist nun auch die für diesen Anfangspunkt gewählte Bezeichnung als „Zentrum der Hauptachsen“ entnommen. Unter *Hauptachsen* sollen nämlich hier, ähnlich wie bei einem Körper im leeren Raume, die Doppelachsen derjenigen stationären Bewegungen des Körpers verstanden werden, welche durch ein bloßes impulsives Kräftepaar zu erzeugen sind, d. h. der Annahme $J = 0$ ($\lambda = \infty$) entsprechen. Die Hauptachsen sind hiernach erste Durchmesser des Körpers, und zwar diejenigen, welche zu den Richtungen der Hauptachsen eines Ellipsoids $G = \text{konst.}$ gehören.

Für die Folge sollen stets die im Körper festen Koordinatenachsen in die Richtungen dreier zueinander senkrechter Hauptachsen gelegt sein, und

setze ich demgemäß für $G(p, q, r)$ einen Ausdruck $\frac{1}{2}(Ap^2 + Bq^2 + Cr^2)$ voraus; die Gleichung (4) für J_1 läßt sich dann in der Form schreiben:

$$(4a) \quad Aup + Bvq + Cwr = JJ_1 - 2F(u, v, w),$$

deren Vorzüge später hervortreten werden.

$\lambda = 0$. Die Annahme $\lambda = 0$ liefert *stationäre Bewegungen ohne Drehung*, also bloße Verschiebungen. Die Richtungen solcher stationären Verschiebungsgeschwindigkeiten sind, wie bereits Kirchhoff angegeben hat, durch die Hauptachsen eines Ellipsoids $E = \text{konst.}$ bezeichnet; die Achsen ihrer Schraubenimpulse sind die zu ihren Richtungen gehörenden zweiten Durchmesser des Körpers. —

Wird ein stationärer Bewegungszustand, wie wir ihn hier betrachten, *stabil* genannt, wenn aus keiner unendlich kleinen Störung des Zustandes im Laufe der Zeit endliche Änderungen seiner Geschwindigkeiten hervorgehen können, so ist eine, durch ein bloßes impulsives Kräftepaar entstandene stationäre Bewegung, etwa eine solche um die der z -Achse parallele Hauptachse — bei der zuletzt getroffenen Wahl des Koordinatensystems — in folgenden Fällen stabil: Wenn das Moment $C < A$ und $< B$, oder $C > A$ und $> B$ ist, mit Ausnahme des Falles, wo $C = A + B$ und nicht zugleich der Schnitt der xy -Ebene mit einer Fläche

$$F(x, y, z) - \frac{1}{4} \left(\frac{\partial^2 F}{\partial x^2} + \frac{\partial^2 F}{\partial y^2} + \frac{\partial^2 F}{\partial z^2} \right) (x^2 + y^2 + z^2) = \text{konst.}$$

eine Ellipse $\frac{x^2}{A} + \frac{y^2}{B} = \text{konst.}$ ist; endlich wenn $A = B = C$ und die z -Achse zugleich eine Hauptachse einer Fläche $F = \text{konst.}$ ist.

Für die übrigen stationären Bewegungen, bei welchen eine impulsive Einzelkraft mitwirkt, kann die Entscheidung über die Stabilität in einfacher Weise von einer gewissen quadratischen Gleichung abhängig gemacht werden. Es erweist sich hier als eine hinreichende Bedingung für Stabilität in dem angegebenen Sinne, wenn bei festgehaltener Größe des Impulses die zweite Variation der gesamten lebendigen Kraft wesentlich positiv ist. Letzteres wieder, und um so mehr Stabilität tritt sicher dann ein, wenn in der Fläche (6), die zu der betrachteten Bewegung gehört, und wo nun die Konstante der rechten Seite *positiv* angenommen sein soll, das reziproke Quadrat desjenigen Durchmessers, mit welchem die Achse der Bewegung parallel ist, numerisch kleiner ist als das reziproke Quadrat irgendeines anderen Durchmessers.

§ 3.

Nunmehr untersuchen wir, welchen Verlauf eine beliebige Bewegung des Körpers mit konstantem Impulse darbietet.

$J = 0$. Ist der Impuls ein bloßes impulsives Kräftepaar, so existiert nur *der* Teil der Bewegung, welcher von den Drehungsgeschwindigkeiten p, q, r abhängt; und letztere haben in jedem Augenblicke dieselben Werte, und ist dadurch auch die Winkelstellung des Körpers im Raume jederzeit dieselbe, als ob der Körper mit gleichem Impulse sich im leeren Raume bewegte, und dabei seine lebendige Kraft um den Trägheitsmittelpunkt den Ausdruck $G(p, q, r)$ hätte. Letztere Bewegung würde nach bekannten Formeln zu berechnen sein, befolgt übrigens auch die weiter unten für den Fall $J > 0$ aufgestellten Sätze, indem aus jenen Sätzen, wenn die Bewegung im leeren Raume vor sich geht, (wenn $E(u, v, w)$ ein Vielfaches von $u^2 + v^2 + w^2$ und F identisch Null ist), die Größe der impulsiven Einzelkraft ohne weiteres herausfällt. Sind die Drehungen des Körpers bereits ermittelt, so findet man die vom Zentrum der Hauptachsen parallel zu irgendeiner im Raume festen Richtung n zurückgelegte Strecke durch das Zeitintegral:

$$- \int \left(\frac{\partial F}{\partial p} \cos(nx) + \frac{\partial F}{\partial q} \cos(ny) + \frac{\partial F}{\partial r} \cos(nz) \right) dt.$$

Dieses Integral bleibt für Richtungen parallel zur Ebene des impulsiven Kräftepaars in der Regel *immer* endlich. Verschwindet die Form F identisch, so übernimmt das Zentrum der Hauptachsen offenbar die Rolle eines festen Punktes.

$J > 0$. Hat der Impuls der Bewegung eine nichtverschwindende impulsive Einzelkraft, so wird durch folgende Gleichungen zunächst ausgedrückt, daß für diese Kraft Größe und *Richtung* im Raume unveränderlich sind:

$$(7) \quad \frac{du}{dt} = rv - qw, \quad \frac{dv}{dt} = pw - ru, \quad \frac{dw}{dt} = qu - pv.$$

Diese Gleichungen sind hinsichtlich p, q, r nicht voneinander unabhängig, sondern liefern die von p, q, r freie Relation $u du + v dv + w dw = 0$, welche zu der Gleichung (3) führt; bringt man sie aber in Verbindung mit der Gleichung (4a) für J_1 , die ebenfalls linear in p, q, r ist, so gelangt man zu Beziehungen:

$$(Au^2 + Bv^2 + Cw^2)p = (JJ_1 - 2F(u, v, w))u + Cw \frac{dv}{dt} - Bv \frac{dw}{dt} \text{ usw.,}$$

mit deren Hilfe man, da $Au^2 + Bv^2 + Cw^2$ hier gewiß von Null verschieden ist, die Drehungsgeschwindigkeiten p, q, r durch die impulsiven Einzelkräfte u, v, w und deren Differentialquotienten nach t darstellen kann. Führt man alsdann für u, v, w solche Funktionen zweier Argumente, e_1 und e_2 , ein, daß die Gleichung $u^2 + v^2 + w^2 = J^2$ identisch erfüllt wird, so spricht sich der Inhalt der nach § 2 für u, v, w, p, q, r bestehenden Differentialgleichungen erster Ordnung, soweit er nicht bereits durch (3),

(4a) und (7) erschöpft ist, in zwei Differentialgleichungen zweiter Ordnung für e_1 und e_2 aus, Gleichungen, die, wie ich nach umständlichen Rechnungen gefunden habe, folgende Auslegung gestatten:

Man fixiere einen beliebigen Punkt und eine beliebige Richtung im Körper; in einem Zeitelement dt sei $d\sigma$ der Weg des Punktes längs der unveränderlichen Achse des Impulses, $d\sigma_1$ der Weg der Richtung um diese Achse; die Projektionen des Punktes auf die in Rede stehende Achse machen die Strecke $d\sigma$ anschaulich, den Winkel $d\sigma_1$ beschreiben die Projektionen der Richtung auf eine zu dieser Achse senkrechte Ebene. *Die Größen e_1 und e_2 sind innerhalb einer beliebigen Zeitperiode solche Funktionen ihrer ersten und ihrer letzten Werte und der Zeit, daß die erste Variation des über die Periode ausgedehnten Integrals*

$$\Phi = \int (T dt - J d\sigma - J_1 d\sigma_1)$$

verschwindet.

Die Integrale $\int dt$, $\int d\sigma$, $\int d\sigma_1$ über die betreffende Periode wollen wir t , σ , σ_1 nennen.

Sind, in bezug auf das im Körper feste Koordinatensystem, a, b, c die Koordinaten des Punktes, α, β, γ die Kosinus der Richtung, so hat man:*)

$$d\sigma = \frac{1}{J} [u(u + cq - br) + v(v + ar - cp) + w(w + bp - aq)] dt,$$

$$d\sigma_1 = J \frac{up + vq + wr - (\alpha u + \beta v + \gamma w)(\alpha p + \beta q + \gamma r)}{u^2 + v^2 + w^2 - (\alpha u + \beta v + \gamma w)^2} dt.$$

Daß es auf die besondere Wahl des Punktes und der Richtung nicht ankommen kann, schließt man aus dem Umstande, daß die Differenz der Gesamtgrößen σ oder σ_1 für zwei Punkte bzw. zwei Richtungen sich ohne Integralzeichen, allein durch die Anfangs- und Endwerte von e_1 und e_2 darstellen läßt.

Nach der Definition von $d\sigma_1$ muß der Winkel σ_1 eine sprungweise Änderung um $\pm \pi$ erleiden, so oft die Richtung, auf welche sich σ_1 bezieht, eine Lage passiert, in welcher sie der Achse des Impulses parallel wird, was offenbar nur in nicht stationären Bewegungen und hier jedesmal nur für einfach unendlich viele Richtungen im Körper eintreten kann. Doch bleibt σ_1 in dem Falle stetig, wenn zu gleicher Zeit die augenblickliche Drehungsachse der Achse des Impulses parallel wird, d. h. $du = 0$, $dv = 0$, $dw = 0$ ist, in welchem Falle nicht auch $d^2u = 0$, $d^2v = 0$, $d^2w = 0$ sein kann, weil sonst die Bewegung stationär weitergehen müßte.

Daß die gesamte lebendige Kraft T konstant ist — wir wollen ihren

*) s. Kirchhoff, Crelles Journal, Bd. 71, S. 255.

konstanten Wert L nennen — drückt sich nach Elimination von p, q, r in der Gleichung aus:

$$(8) \quad E(u, v, w) + \frac{1}{2} \frac{[JJ_1 - 2F(u, v, w)]^2}{Au^2 + Bv^2 + Cw^2} + \frac{1}{2} \frac{BC \left(\frac{dw}{dt}\right)^2 + CA \left(\frac{dv}{dt}\right)^2 + AB \left(\frac{du}{dt}\right)^2}{Au^2 + Bv^2 + Cw^2} = L,$$

die offenbar dazu verwandt werden kann, das Element dt aus den Differentialgleichungen für e_1 und e_2 herauszuschaffen. Man kommt dadurch zu folgendem rein geometrischen Resultate:

Ist die Arbeit des Impulses der Bewegung (L) und seine Größe (J, J_1) bekannt, und sind von den Stellungen, welche die Richtung der im Raume unveränderlichen Achse des Impulses zu den verschiedenen Zeiten in bezug auf den Körper einnimmt, irgend zwei ($e_1^0, e_2^0; e_1, e_2$) gegeben, so ist

$$\Psi = \int_{(e_1^0, e_2^0)}^{(e_1, e_2)} (2Ldt - Jd\sigma - J_1d\sigma_1)$$

oder

$$2Lt - J\sigma - J_1\sigma_1$$

auf dem wirklichen Wege dieser Richtung im Körper ein Minimum (beziehlich, wenn die zwei Stellungen weiter auseinander liegen, ein Grenzwert).

Es empfiehlt sich, für die Bestimmungsstücke e_1 und e_2 dieser Stellungen die zwei elliptischen Koordinaten zu nehmen, die einem Punkte, dessen rechtwinklige Koordinaten (abgesehen von der Dimension)

$$x = \frac{1}{\sqrt{A}} \frac{u}{J}, \quad y = \frac{1}{\sqrt{B}} \frac{v}{J}, \quad z = \frac{1}{\sqrt{C}} \frac{w}{J}$$

sind, auf dem, mit dem Körper fest verbundenen Ellipsoide

$$Ax^2 + By^2 + Cz^2 = 1$$

zukommen. Die Gleichung (8) zeigt nämlich, daß die Geschwindigkeit, die ein so gewählter Punkt in seiner Bahn auf diesem Ellipsoide hat, allein eine Funktion seines Ortes auf letzterem ist; und heißt diese Geschwindigkeit $\frac{ds}{dt}$ oder s' , so nimmt dazu das Element des Integrals Ψ einen Ausdruck an:

$$\frac{ABC}{A^2x^2 + B^2y^2 + C^2z^2} s' ds + F_1 de_1 + F_2 de_2.$$

Hierin sind die Funktionen F_1 und F_2 Null (für $a, b, c = 0, 0, 0$), wenn die Form F identisch verschwindet und zugleich die Konstante J_1 Null ist. —

Auf die im vorstehenden entwickelten Minimalsätze sind nun die Methoden von Hamilton und Jacobi sehr einfach anzuwenden.*) Drückt man in $T - J \frac{d\sigma}{dt} - J_1 \frac{d\sigma_1}{dt}$ die Größen u, v, w, p, q, r durch

*) Siehe C. G. J. Jacobi, Vorlesungen über Dynamik, Vorl. XIX und XXII.

$$e_1, e_2, \frac{de_1}{dt} = e_1', \frac{de_2}{dt} = e_2', J, J_1$$

aus, bezeichnet die entstehende Funktion mit Φ , führt $\frac{\partial \Phi}{\partial e_1'} = f_1, \frac{\partial \Phi}{\partial e_2'} = f_2$ ein, und stellt $f_1 e_1' + f_2 e_2' - \Phi$ als Funktion $\psi(e_1, e_2, f_1, f_2, J, J_1)$ dar, so lauten die Differentialgleichungen für e_1, e_2, f_1, f_2 :

$$dt : de_1 : de_2 : df_1 : df_2 = 1 : \frac{\partial \psi}{\partial f_1} : \frac{\partial \psi}{\partial f_2} : -\frac{\partial \psi}{\partial e_1} : -\frac{\partial \psi}{\partial e_2}.$$

Die Gleichung $T = L$ geht in $\psi = L$ über und wird durch Einsetzung von $\frac{\partial \psi}{\partial e_1}, \frac{\partial \psi}{\partial e_2}$ für f_1, f_2 zu einer partiellen Differentialgleichung, welcher das Integral Ψ genügt, wenn man dasselbe als Funktion der Werte von e_1, e_2 , an seiner oberen Grenze ansieht, die Werte dieser Größen für die untere Grenze, e_1^0, e_2^0 , dagegen als konstant betrachtet.

Sobald von dieser partiellen Differentialgleichung eine Lösung gefunden ist, welche eine willkürliche Konstante, außer der bloß additiv hinzutretenden, enthält, können unmittelbar sämtliche Integralgleichungen der Bewegung des Körpers hingeschrieben werden. Bezeichnet man mit Ψ die Differenz aus einer solchen Lösung und dem Werte, den die Lösung für das System $e_1 = e_1^0, e_2 = e_2^0$ annimmt, und mit M jene, noch in dieser Differenz vorkommende Konstante, so sind

$$(9) \quad \frac{\partial \Psi}{\partial L} = t, \quad \frac{\partial \Psi}{\partial M} = 0$$

die Gleichungen, welche e_1 und e_2 als Funktionen von t definieren; eine eingehendere Betrachtung des Ausdrucks von Ψ als Integral führt zu:

$$\sigma_1 = -\frac{\partial \Psi}{\partial J_1},$$

und endlich ist:

$$\sigma = \frac{1}{J} (-\Psi + 2Lt - J_1 \sigma_1).$$

Diese Gleichungen bestimmen nun zu jeder Zeit den Ort des Körpers vollständig. Denn man erhält ein im Raume festes rechtwinkliges Koordinatensystem ξ, η, ζ , indem man folgende Annahmen macht: die ζ -Achse soll mit der Achse des Impulses identisch sein, der Anfangspunkt mit dem Ausgangspunkte der Strecke σ , die Richtung der ξ -Achse mit der Ausgangsrichtung des Winkels σ_1 , und endlich soll vermöge der η -Achse das Koordinatensystem der ξ, η, ζ dem im Körper festen Systeme der x, y, z kongruent sein. Die so definierten Achsen ξ, η, ζ kann man aber in jedem Augenblicke vom Körper aus mit Hilfe der letzten Gleichungen, nämlich durch die Werte von e_1, e_2, e_1', e_2' ; σ, σ_1 wiederfinden. Durch die vier ersten Größen oder vielmehr durch u, v, w, p, q, r findet man nach (5) die ζ -Achse, dann durch σ den Anfangspunkt, durch σ_1 die Richtung der ξ -Achse.

Für die Entfernung, die irgendein Punkt des Körpers von der Achse des Impulses hat, gewinnt man leicht, durch Berücksichtigung des Umstandes, daß die lebendige Kraft T eine wesentlich positive Form ist, eine obere Grenze von der Art: $\frac{\sqrt{L}}{J}$, multipliziert in eine von den Koeffizienten der Form T abhängige Konstante, so daß diese Entfernung immer endlich bleibt. —

A. Clebsch*) hat bemerkt, daß für die von Kirchhoff aufgestellten Differentialgleichungen der Multiplikator eine Konstante wird, und er hat daraus den Schluß gezogen, daß eine vollständige Integration dieser Gleichungen durch Quadraturen gelingt, sowie den drei Integralen mit den Konstanten J, J_1, L irgendein, gleichfalls von t freies, viertes Integral hinzugefügt werden kann. Weit mehr, als die bloße Anwendung des Prinzips des letzten Multiplikators ergibt, leisten nach der hier vorgenommenen Transformation jener Gleichungen die Sätze von Jacobi über diejenigen Probleme der Mechanik, welche nur zwei zu bestimmende Größen enthalten. Bringt man das vierte Integral auf eine der Gleichung $\psi = L$ analoge Form:

$$\chi(e_1, e_2, f_1, f_2, J, J_1) = \text{konst.} = M,$$

und ermittelt aus diesen zwei Gleichungen f_1 und f_2 als Funktionen von e_1 und e_2 , so ist nach jenen Sätzen $f_1 de_1 + f_2 de_2$ ein vollständiges Differential und

$$(10) \quad \Psi = \int_{(e_1^0, e_2^0)}^{(e_1, e_2)} (f_1 de_1 + f_2 de_2),$$

woraus die Quadraturen für alle Integralgleichungen zu ersehen sind.

Indem Clebsch den Ansatz machte, daß das vierte Integral eine ganze homogene Funktion ersten oder zweiten Grades der Geschwindigkeiten des Körpers sein sollte, ergaben sich ihm im ganzen zwei Fälle, in welchen ein derartiges viertes Integral existiert; in unserer Bezeichnung sind dieselben sehr einfach zu charakterisieren:

Wenn die drei Flächen $E = \text{konst.}$, $F = \text{konst.}$, $G = \text{konst.}$ Rotationsflächen um eine gemeinsame Achse sind, so ist die Drehungsgeschwindigkeit um diese Achse konstant. Von der Art, wie die Bewegung alsdann vonstatten geht, hat neuerdings Herr Halphen**) ein sehr anschauliches Bild auf Grund von Eigenschaften elliptischer Funktionen entworfen. Kommt der Umstand hinzu, daß die Form F identisch verschwindet, so sind Rotationskörper — nach Gestalt und Massenverteilung — denkbar, die in

*) Mathematische Annalen, Bd. 3, S. 238—262.

**) Journal de Liouville, 1888. 4^e série, T. IV, pp. 1—81. — Auch in dem *Traité des fonctions elliptiques et de leurs applications*, T. II.

ihren Bewegungen ohne Einfluß von Kräften vollständig mit dem gegebenen Körper übereinstimmen, ein Fall, für welchen bereits Kirchhoff die Möglichkeit einer Integration durch Quadraturen gezeigt hatte.

Wenn ferner F identisch Null ist, und die durch zwei Flächen $E = \text{konst.}$ und $G = \text{konst.}$ bestimmte Flächenschar eine Kugel aufweist, d. h. eine lineare Relation

$$(11) \quad 2E(x, y, z) + l(Ax^2 + By^2 + Cz^2) = m(x^2 + y^2 + z^2)$$

besteht, so ergibt sich als ein viertes Integral:

$$l(BCu^2 + CAv^2 + ABw^2) + A^2p^2 + B^2q^2 + C^2r^2 = \text{konst.} = 2M.$$

Ist $G = \text{konst.}$ die betreffende Kugel, also $l = \infty$, $A = B = C = \frac{m}{l}$, so scheint dieses Integral nur in die Gleichung (3) überzugehen, weshalb auch Clebsch diesen Fall besonders untersucht hat; indessen erweist sich die lineare Verbindung $lM - mL - \left(\frac{1}{2}(A + B + C)lm - m^2\right)J^2$ aus den drei Gleichungen für M , L und J^2 auch hier als ein neues Integral, nachdem man aus derselben zuerst mit Hilfe von (11) A, B, C eliminiert und dann für $\frac{m}{l}$ den gemeinsamen Wert dieser Größen und $l = \infty$ gesetzt hat; von der Möglichkeit, daß sowohl $G = \text{konst.}$ wie $E = \text{konst.}$ Kugeln sind, sehen wir dabei ab, dann würde übrigens jede Bewegung eine stationäre sein. In diesem zweiten Falle, der durch Quadraturen zu erledigen wäre, stieß Clebsch auf Schwierigkeiten bei dem Versuche, die Bewegung als explizite Funktion der Zeit darzustellen. Dieses Ziel hat dann Herr H. Weber*) durch Benutzung von Thetareihen mit zwei Argumenten für den Fall vollständig erreicht, wo die Konstante J_1 Null, der Impuls der Bewegung also eine einzelne impulsive Kraft ist. Hier wird durch die Gleichungen (9) und (10) für jeden Fall dasjenige Umkehrproblem in einfachster Weise bezeichnet, dessen Lösung den Kernpunkt bei der eigentlichen Berechnung der Bewegung bildet.

*) Mathematische Annalen, Bd. 14, S. 173–206.

XXIX.

Kapillarität.

(Enzyklopädie der mathematischen Wissenschaften. V 1. Heft 4. S. 558—613.)

Literatur.

Lehrbücher und Monographien.

- Thomas Young, Essay on the cohesion of fluids, London Philosophical Transactions of the Royal Society, 1805.
- P. S. Laplace, Théorie de l'action capillaire, Supplément au livre dixième de la Mécanique céleste, Paris 1806; Supplément à la théorie de l'action capillaire, Paris 1807; letzteres auch in Blainville, Journal de Physique, T. 2 (1806), p. 120; beides zusammen in Oeuvres, T. 4, Paris 1845 und Oeuvres, T. 4, Paris 1880; deutsch in Annalen der Physik, Bd. 33 (1809).
- C. Fr. Gauss, Principia generalia theoriae figurae fluidorum in statu aequilibrii, Commentationes societatis regiae scientiarum Gottingensis recentiores, vol. 7 (1830), wiederabgedruckt in Werke, Bd. 5, S. 29, Göttingen 1867, übersetzt von R. H. Weber in Ostwalds Klassiker der exakten Wissenschaften, Nr. 135, Leipzig 1903.
- S. D. Poisson, Nouvelle théorie de l'action capillaire, Paris 1831; in einem kurzen Auszuge übersetzt von Link, in Annalen der Physik und Chemie, Bd. 25 (1832), S. 270; Bd. 27 (1833), S. 193.
- E. Betti, Teoria della capillarità, Nuovo Cimento, T. 25 (1867), p. 81, 225.
- J. Plateau, Statique expérimentale et théorique des liquides soumis aux seules forces moléculaires, 2 Bde., Gand 1873.
- J. C. Maxwell, Capillary action, Encyclopaedia Britannica 1875 (9. Aufl.), wiederabgedruckt in Scientific papers, vol. 2, Cambridge 1890, p. 541.
- J. W. Gibbs, Equilibrium of heterogeneous substances, Connecticut Academical Transactions, vol. 3, New Haven 1876 und 1878, deutsch in: Thermodynamische Studien, übersetzt von W. Ostwald, Leipzig 1892, S. 66.
- É. Mathieu, Théorie de la capillarité, Paris 1883.
- J. D. van der Waals, Die Continuität des gasförmigen und flüssigen Zustandes, Leipzig 1881, 2. Aufl. 1. Teil, Leipzig 1899.
- B. Weinstein, Untersuchungen über Capillarität, Annalen der Physik und Chemie, Bd. 27 (1886), S. 544.
- J. D. van der Waals, Thermodynamische Theorie der Kapillarität unter Voraussetzung stetiger Dichteänderung, Verhandelingen der Koninklijke Akademie van wetenschappen, Amsterdam, Deel I, Nr. 8 (1893), übersetzt in Zeitschrift für physikalische Chemie, Bd. 13 (1894), S. 657; Archives Néerlandaises des Sciences exactes et naturelles; T. 28 (1894).
- Fr. Neumann, Vorlesungen über die Theorie der Capillarität, herausg. von A. Wangerin, Leipzig 1894.
- H. Poincaré, Capillarité (Leçons), Paris 1895.
- Literaturnachweis: Domke, Wissenschaftliche Abhandlungen der K. Normaleichungskommission, Berlin 1902, S. 81—99.

1. Kapillarität und Kohäsion.

An den Grenzen einer Flüssigkeit gegen die sie umgebenden Medien äußert sich eine besondere Form der Energie, welche namentlich Gestalt und Lage der freien Grenzflächen beeinflusst, wie z. B. die Höhe des Ansteigens der Flüssigkeit in einer Kapillarröhre, und welche von diesem speziellen Phänomene her allgemein den Namen *Kapillarität* erhalten hat. Diese Energie ist sichtlich verknüpft mit dem, sei es nun diskontinuierlichen, sei es schnellen kontinuierlichen Wechsel des Zustandes über eine Trennungsfläche hinüber. Theoretisch wird sie entweder unmittelbar angesetzt mit einem Term, welcher dem Flächeninhalt der Trennungsfläche proportional ist, und wobei der Proportionalitätsfaktor eine *Spannung in der Oberfläche* angibt. Oder aber sie wird erklärt als die potentielle Energie von besonderen Anziehungskräften, welche zwischen allen Massenteilchen untereinander, jedoch mit wahrnehmbarem Betrage nur auf äußerst kleine Distanzen wirksam sind. Alsdann ist ein überwiegender Anteil dieser Energie dem Volumen der Flüssigkeit proportional, wobei der Proportionalitätsfaktor, negativ genommen, einen *Druck im Raume* angibt, den man die *Kohäsion* der Flüssigkeit nennt.

Es soll hier die erstere Auffassung der Kapillarität vorangestellt werden, welche dieser Energieform die Trennungsflächen als ausschließlichen Sitz zuweist. Diese Auffassung erscheint hernach als ein mathematisch einfacher Grenzfall der anderen tieferen Auffassung, welche die ganzen Massen als Spielraum von Kohäsionskräften annimmt.

I. Kapillarität als Flächenenergie.

2. Oberflächenenergie und deren Variation.

Eine Trennungsfläche zwischen einer Flüssigkeit A und einem zweiten Medium B ist verknüpft mit einer potentiellen Energie $T_{AB}F_{AB}$, wobei F_{AB} den Flächeninhalt der Fläche und T_{AB} eine von den beiderseits angrenzenden Medien abhängende Konstante ist. Dabei sind A und B homogen gedacht und sollen Änderungen von Temperatur und Dichte zunächst nicht in Betracht kommen. T_{AB} hat die Dimensionen $\frac{\text{erg}}{\text{cm}^2}$ und heißt *Oberflächenspannung* von A gegen B . Unter der Oberflächenspannung schlechthin versteht man für eine Flüssigkeit diejenige gegen ihren gesättigten Dampf, wovon die einer Trennungsfläche gegen Luft*) vielfach

*) Von einer Oberflächenspannung gegen eine gasförmige Phase kann, streng genommen, nur bei Flüssigkeiten die Rede sein, welche mit der gasförmigen Phase in chemischem Gleichgewicht koexistieren.

keine Verschiedenheit zeigt. Für Wasser gegen Luft ist

$$T = 74 \frac{\text{erg}}{\text{cm}^2} = 0,075 \frac{\text{gr Gewicht}}{\text{cm}},$$

für Quecksilber gegen Luft $T = 0,55$ gr Gewicht/cm, für Quecksilber gegen Wasser $T = 0,42$ gr Gewicht/cm.

Die Folgerungen aus dem Bestehen des Terms $T_{AB}F_{AB}$ in der Energie ließen sich am kürzesten darlegen auf Grund der Bemerkung, daß genau derselbe Ausdruck der Energie Platz greifen würde für eine unendlich dünne elastische Haut, welche die Trennungsfläche bedeckt, wenn in ihr überall eine konstante Spannung $= T_{AB}$ herrscht, d. h. jede in ihr angebrachte Schnittlinie an beiden Ufern einen von dem anderen fort gerichteten Zug T_{AB} auf die Längeneinheit erfährt. Diesen Vergleich von vornherein einzuführen, hieße aber, die Schwierigkeit, welche für die Theorie der Kapillarität in der Notwendigkeit der Annahme von Druckdiskontinuitäten transversal zu einer Trennungsfläche liegt, nicht beseitigen, sondern nur sie einem anderen Kapitel der Mechanik zuschieben. Um hinsichtlich der Voraussetzungen der Theorie Klarheit zu gewinnen, ist es deshalb notwendig, im wesentlichen den Gang von Gauß einzuhalten und den Einfluß des Terms TF der Energie auf Grund eines allgemeinen Prinzips der Mechanik zu verfolgen, wie des Prinzips, daß Gleichgewicht durch ein Minimum der potentiellen Energie oder, bei Berücksichtigung auch thermodynamischer Umstände, durch ein Minimum der Energie bei konstanter Entropie charakterisiert wird.

Hiernach muß vor allem die virtuelle Veränderlichkeit von TF betrachtet werden. Gauß hat, indem er bei diesem Anlasse überhaupt die Prinzipien für die Variation von Doppelintegralen mit veränderlichen Grenzen schuf, eine fundamentale Transformation für die Variation von TF entwickelt. Man kann sich allerdings, worauf auch Gauß beiläufig hinweist, das Resultat dieser Transformation durch infinitesimale Betrachtungen leicht plausibel machen, indem man eine beliebige unendlich kleine Verrückung einer Fläche in eine erste Verrückung, wobei jeder Punkt normal zur Fläche fortschreitet, und eine zweite Verrückung, wobei jeder Punkt tangential zur Fläche fortschreitet, zerlegt. Doch erscheint es angemessen, hier auch das eigentliche analytische Prinzip jener Umwandlung, welches in einer gewissen partiellen Integration besteht, anzugeben. In Anbetracht des Umstandes jedoch, daß die Variationsrechnung, namentlich in Hinsicht auf Probleme mit Nebenbedingungen, wie sie hier schließlich vorliegen werden, noch keine allgemein anerkannte Darstellung besitzt, auf die sonst einfach zu verweisen wäre, mag es zur Durchsichtigkeit beitragen, wenn wir nur auf bekanntere Hilfsmittel der Integralrechnung rekurrieren.

Die Trennungsfläche F_{AB} , die wir uns berandet denken wollen, durch-

laufe bei einer virtuellen Bewegung, während welcher der Parameter w von $w = 0$ an wächst, eine Schar von Flächen $F(w)$. Wir konstruieren auf jeder Fläche die zwei zueinander orthogonalen Scharen von Krümmungslinien und sodann zwei Flächenscharen $u_2 = \text{konst.}$, $u_1 = \text{konst.}$, welche aus den $F(w)$ gerade diese Krümmungslinien heraus schneiden. Wir stellen uns der Einfachheit halber die Flächen $F(w)$ sämtlich als auseinanderliegend und derart vor, daß in dem von ihnen erfüllten Gebiete jedem Punkte sich hiernach Werte u_1 , u_2 , w eindeutig zuordnen. Die Richtungen von einem Punkte u_1 , u_2 , w aus, in denen nur u_1 , nur u_2 , nur w und zwar zunehmend variiert, sollen in Argumenten trigonometrischer Funktionen kurz durch u_1 , u_2 , w angedeutet werden, und n bedeute die nach B hin gerichtete Normale auf $F(w)$; die u_1 , u_2 , n -Richtung sollen stets ein Rechtsschraubensystem wie die Koordinatenachsen x , y , z bilden (Fig. 1).

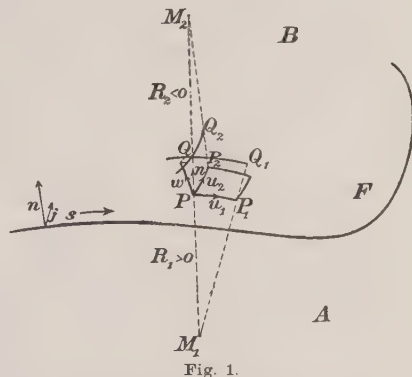


Fig. 1.

Die Berandung von $F(w)$ wird durch eine Gleichung $u_2 = \chi(u_1, w)$ dargestellt sein, und da wir irgendeine Funktion von u_1 und w als einen neuen Parameter an Stelle von u_1 einführen können, so dürfen wir diese Gleichung als w nicht enthaltend: $u_2 = \chi(u_1)$ annehmen. Das Quadrat des Linienelements in dem von den $F(w)$ erfüllten Gebiete wird sich

$$(1) \quad ds^2 = dx^2 + dy^2 + dz^2 \\ = L_1^2 (du_1 - l_1 dw)^2 + L_2^2 (du_2 - l_2 dw)^2 + N^2 dw^2$$

schreiben lassen, wobei $L_1 > 0$, $L_2 > 0$ und $\frac{N}{\cos(wn)} = L > 0$ sei, also N mit dem Vorzeichen von $\cos(wn)$ gerechnet werde. Nun ist

$$F = \iint L_1 L_2 du_1 du_2, \\ (2) \quad \frac{dF}{dw} = \iint \left(L_1 \frac{dL_2}{dw} + L_2 \frac{dL_1}{dw} \right) du_1 du_2,$$

wo das Doppelintegral sich über das Innere von $u_2 = \chi(u_1)$ erstreckt.

Der Ausdruck hier gestattet auf Grund der charakteristischen Eigenschaft der Krümmungskurven, daß die Normalen längs ihnen eine abwickelbare Fläche bilden, eine wichtige Umformung durch Produktintegration*). Zu einem Punkte $P(u_1, u_2, w)$ liegt auf der benachbarten

*) Die zunächst folgende infinitesimale Betrachtung dient nur dazu, diese Eigenschaft der Krümmungslinien schnell in eine Formel umzusetzen.

Fläche $F(w + dw)$ in der kleinstmöglichen Distanz $Ndw = dn$, also auf der Normalen von $F(w)$ der Punkt $Q(u_1 + l_1 dw, u_2 + l_2 dw, w + dw)$. Entsprechend liege normal über $P_1(u_1 + du_1, u_2, w)$ auf $F(w + dw)$ der Punkt Q_1 ; es sei M_1 der Krümmungsmittelpunkt der Krümmungslinie PP_1 auf $F(w)$, also der Treffpunkt der Geraden PQ und P_1Q_1 , R_1 der Krümmungsradius M_1P und zwar positiv, falls M_1P die Richtung n nach B hin hat, anderenfalls negativ. Man hat $PP_1 = L_1 du_1$. Aus der Ähnlichkeit der Dreiecke M_1PP_1 , M_1Q_1 und im Hinblick auf (1) folgt

$$\frac{PQ}{M_1P} = \frac{Q_1Q_1 - PP_1}{PP_1}, \quad \frac{dn}{R_1} = \frac{1}{L_1} \left(\frac{\partial L_1}{\partial n} dn + L_1 \frac{\partial l_1}{\partial u_1} dw \right),$$

d. i. nach Fortlassung des Faktors dw

$$\frac{N}{R_1} = \frac{1}{L_1} \left(\frac{\partial L_1}{\partial u_1} l_1 + \frac{\partial L_1}{\partial u_2} l_2 + \frac{\partial L_1}{\partial w} + L_1 \frac{\partial l_1}{\partial u_1} \right).$$

Eine entsprechende Relation gilt für den zweiten Hauptkrümmungsradius R_2 in P und durch Addition beider entsteht

$$N \left(\frac{1}{R_1} + \frac{1}{R_2} \right) L_1 L_2 = L_1 \frac{\partial L_2}{\partial w} + L_2 \frac{\partial L_1}{\partial w} + \frac{\partial L_1 L_2 l_1}{\partial u_1} + \frac{\partial L_1 L_2 l_2}{\partial u_2}.$$

Macht man hiervon in (2) Gebrauch und führt partielle Integrationen nach u_1 und u_2 aus, so ergibt sich

$$\frac{dF}{dw} = \int_{(f)} N \left(\frac{1}{R_1} + \frac{1}{R_2} \right) df - \int_{(s)} L_1 L_2 \left(l_1 \frac{du_2}{ds} - l_2 \frac{du_1}{ds} \right) ds,$$

darin bedeutet allgemein df das Flächenelement, ds das Randlinienelement von $F(w)$ in positivem Umlauf um die Normale n . Bezeichnet man mit j die Richtung, die normal auf ds ins Innere der Fläche $F(w)$ hineinführt (s. Fig. 1), so ist

$$L_1 \frac{du_1}{ds} = \cos(u_2 j), \quad L_2 \frac{du_2}{ds} = -\cos(u_1 j), \\ -L_1 l_1 = L \cos(u_1 w), \quad -L_2 l_2 = L \cos(u_2 w);$$

schreibt man noch $L \cos(wj) dw = dj$, so entsteht daher aus der letzten Gleichung die folgende Darstellung der ersten Ableitung der Kapillarenergie TF nach dem Variationsparameter w :

$$(3) \quad T \frac{dF}{dw} = T \int_{(f)} \frac{dn}{dw} \left(\frac{1}{R_1} + \frac{1}{R_2} \right) df - T \int_{(s)} \frac{dj}{dw} ds,$$

die insbesondere für $w = 0$ anzuwenden sein wird. Darin bedeuten dann die $dn = Ndw$ für die Punkte der Trennungsfläche die zur Fläche normal und die $dj = L \cos(wj) dw$ für die Punkte ihres Randes die in die Fläche fallenden, zum Rande normalen Komponenten der einem Zuwachs dw entsprechenden Verrückungen; hierauf fußend kann man sich die Transformation (3), wie schon oben angedeutet, unmittelbar geometrisch plau-

sibel machen.*) $\frac{1}{2} \left(\frac{1}{R_1} + \frac{1}{R_2} \right)$ soll, wie die Vorzeichen von R_1 und R_2 oben festgelegt sind, die *mittlere Krümmung* der Stelle df nach B hin heißen.

Da in (3) die Parameter der Krümmungskurven wieder eliminiert sind, so ist diese Darstellung nicht an die anfänglich betreffs der Koordinaten u_1, u_2, w gemachte Beschränkung gebunden.

Die Volumina von A und B seien V_A, V_B , ihre Dichten ϱ_A, ϱ_B . Wir merken noch an, daß bei den fraglichen virtuellen Verrückungen die Volumina gemäß

$$(4) \quad \frac{dV_A}{dw} = \int_{F_{AB}} \frac{dn}{dw} df, \quad \frac{dV_B}{dw} = - \int_{F_{AB}} \frac{dn}{dw} df$$

variieren, ferner die potentielle Energie bezüglich der Schwerkraft für beide Medien zusammen die Ableitung nach w :

$$(5) \quad g(\varrho_A - \varrho_B) \int z \frac{dn}{dw} df$$

erfährt; dabei ist die z -Achse *vertikal nach oben* gedacht.

3. Differentialgleichung für eine freie Oberfläche.

Die Trennungsfläche von A gegen B sei frei beweglich (B eine Flüssigkeit wie A oder ein Gas), und neben der Kapillarität komme nur noch die Schwere in Betracht. Stabiles Gleichgewicht des Systems wird durch ein Minimum der potentiellen Energie gegenüber allen virtuellen Verrückungen charakterisiert. Nun liegt aber eine Nebenbedingung in der Konstanz des Gesamtvolumens von A (oder von B , vgl. (4)) vor. Um dieser Nebenbedingung Rechnung zu tragen, ziehen wir die Regeln der Differentialrechnung für ein sogenanntes relatives Extremum heran.

Wir denken uns wieder die Trennungsfläche F_{AB} als das Element $w = 0$ einer beliebigen von einem Parameter w abhängenden Schar von Flächen $z = \psi(x, y, w)$, welche alle den Rand gemein haben mögen. Der Nebenbedingung würde allerdings, während w sich verändert, nicht mehr genügt werden. Denken wir uns aber noch eine beliebige zweite solche Schar von Flächen $z = \psi^*(x, y, w^*)$, welche wieder für $w^* = 0$ von der gegebenen Fläche ausgeht, und erweitern wir diese zwei einparametrischen Scharen irgendwie zu einer Flächenschar mit zwei Parametern $z = \psi(x, y, w, w^*)$, welche für $w^* = 0$ in die erste, für $w = 0$ in die zweite einparametrische Schar übergeht, so wird die Größe des Volumens V_A bei den Flächen dieser

*) Wir haben im übrigen die gebräuchlichen Zeichen $\delta F, \delta w$ usw. der ersten Variationen vermieden, um evident zu machen, daß es sich schließlich nur um Differentialquotienten im gewöhnlichen Sinne handelt.

allgemeineren Schar eine Funktion $V_A(w, w^*)$ der zwei Parameter sein, und innerhalb der zweiparametrischen Schar gibt uns diejenige einparametrische Schar, welche durch die Bedingung $V_A(w, w^*) = V_A(0, 0)$ ausgeschieden wird, jetzt eine tatsächliche virtuelle Bewegung der Trennungsfläche. Danach haben wir die Bedingung zu formulieren, daß unter allen Flächen der zweiparametrischen Schar, für welche $V_A(w, w^*) = V_A(0, 0)$ ist, die Fläche $w = 0, w^* = 0$ das Minimum der potentiellen Energie

$$E = T_{AB} F_{AB} + g \varrho_A \int_A z dv + g \varrho_B \int_B z dv$$

liefert; (die Bezeichnung hier ist so zu verstehen, daß dv in dem ersten Integral die Volumenelemente von A , in dem zweiten diejenigen von B durchläuft). Für dieses Extremum mit einer Nebenbedingung liefert nun die Differentialrechnung in bekannter Weise die zwei Gleichungen

$$\frac{\partial E}{\partial w} + \lambda_{AB} \frac{\partial V_A}{\partial w} = 0, \quad \frac{\partial E}{\partial w^*} + \lambda_{AB} \frac{\partial V_A}{\partial w^*} = 0 \quad (w = 0, w^* = 0)$$

mit einer geeigneten Konstante λ_{AB} . Die zweite Gleichung dient uns jetzt nur dazu, um zu erkennen, daß der Wert von λ_{AB} in keiner Weise von der beliebig angenommenen ersten Schar $z = \psi(x, y, w)$ abhängt, also für die Trennungsfläche F_{AB} an sich eine bestimmte Bedeutung hat; und bei ausdrücklicher Hinzunahme dieser Tatsache vertritt die erste Gleichung bereits das System der beiden. Danach muß (vgl. (3), (4), (5)) mit einer geeigneten Konstante λ_{AB} , die sich schließlich aus dem Werte von V_A bestimmen wird, die Bedingung

$$\int_{F_{AB}} N(T_{AB} \left(\frac{1}{R_1} + \frac{1}{R_2} \right) + g(\varrho_A - \varrho_B)z + \lambda_{AB}) df = 0$$

gelten. Dabei unterliegt vermöge der willkürlichen Wahl der Schar $F(w)$ die Funktion $N = \frac{dn}{dw}$ auf der Fläche F_{AB} einzig der Beschränkung, daß sie durchweg stetig ist und am Rande gleich Null genommen wird. Für N in diesem Umfange kann das vorstehende Integral nur dann beständig gleich Null ausfallen, wenn der Faktor von N an jeder Stelle innerhalb F_{AB} verschwindet, d. h. die Gestalt der freien Oberfläche muß der Differentialgleichung

$$(6) \quad T_{AB} \left(\frac{1}{R_1} + \frac{1}{R_2} \right) + g(\varrho_A - \varrho_B)z + \lambda_{AB} = 0$$

genügen.*)

Es kann F_{AB} auch aus mehreren getrennten Stücken bestehen und λ_{AB} hat für die verschiedenen Stücke denselben Wert.

*) Laplace, Supplément au livre X de la Mécanique céleste, no. 4.

Ein Minimum von E ist hier jedenfalls nur möglich, wenn $T_{AB} \geq 0$ ist, da sonst durch ein Hin- und Herfalten der Trennungsfläche an ihrem Orte sich E beliebig verringern ließe.

Es möge $\varrho_A \neq \varrho_B$ sein. Setzen wir

$$z_{AB} = -\frac{\lambda_{AB}}{g(\varrho_A - \varrho_B)},$$

so wird durch $z = z_{AB}$ eine bestimmte horizontale Ebene angewiesen, welche *Niveauebene* heißen soll. Verlegen wir $z = 0$ nach der Niveauebene, so folgt die Gleichung (6) in der Gestalt

$$(6a) \quad T_{AB} \left(\frac{1}{R_1} + \frac{1}{R_2} \right) = -g(\varrho_A - \varrho_B)z.$$

Danach ist an jeder Stelle der Trennungsfläche aus der mittleren Krümmung nach B hin sofort auf die Lage der Niveauebene zu schließen. Ist $\varrho_A > \varrho_B$, so liegen die Stellen der Trennungsfläche, wo diese mittlere Krümmung positiv ist, d. h. die Fläche nach B hin konvex-konvex oder konvex-konkav mit größerem Betrage der ersteren Hauptkrümmung ist, unterhalb der Niveauebene und zwar um so tiefer darunter, je stärker jene mittlere Krümmung ist; Stellen, wo diese Krümmung nach B hin negativ ist, liegen oberhalb der Niveauebene, Stellen, wo sie Null ist, notwendig genau in Höhe dieser Ebene (Fig. 2). Insbesondere kann die Trennungsfläche asymptotisch eben nur in Höhe dieser Ebene sich gestalten, wodurch ihre Bezeichnung als Niveauebene begründet ist.

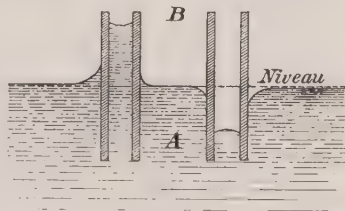


Fig. 2.

4. Randwinkel.

Zur vollständigen Festlegung von F_{AB} sind außer der Differentialgleichung (6) weitere Bedingungen für den Rand der Fläche dort, wo A und B an dritte Medien C grenzen, erforderlich.

Grenzen drei Flüssigkeiten A, B, C mit drei Trennungsflächen F_{AB}, F_{AC}, F_{BC} , in denen die Oberflächenspannungen T_{AB}, T_{AC}, T_{BC} herrschen, längs einer Kurve zusammen (Fig. 3), so würde zwar mit jedem virtuellen Bewegungszustand der Randkurve stets auch der entgegengesetzte Bewegungszustand für sie virtuell sein; immerhin wollen wir (weil hernach auch ein Fall von nicht in diesem Sinne



Fig. 3.

umkehrbaren Verrückungen in Betracht kommt), zunächst nur von dem vorausgesetzten Gleichgewichtszustand an (nicht durch ihn hindurch)

variieren, und wir denken uns eine von einem Parameter $w (\geq 0)$ abhängende Schar von Lagen dieser Randkurve, mit den nämlichen Endpunkten, falls nicht die Kurve geschlossen ist, und von gewissen damit verbundenen Lagen der drei Trennungsf lächen; dabei sollen diese stets gemeinsam an der Randkurve ansetzen, ihre sonstigen Randteile aber fest behalten und zugleich in ihren inneren Partien solche Deformationen erfahren, daß die ganzen Volumina von A, B, C unverändert bleiben. Um zum Ausdruck zu bringen, daß die Gesamtenergie E im Gleichgewichtszustand für $w = 0$ am kleinsten ist, haben wir dann in Anbetracht der erst einseitigen Variation bloß die Ungleichung

$$\frac{dE}{dw} \geq 0 \quad (w = 0)$$

zu fordern. Da aber für die Trennungsf lächen gemäß (6) bereits solche Gleichungen feststehen, daß die hierin als Flächenintegrale auftretenden Terme für sich Null sind, so zieht sich diese Ungleichung gemäß (3) zu

$$-\int L(T_{AB} \cos(wj_{AB}) + T_{AC} \cos(wj_{AC}) + T_{BC} \cos(wj_{BC})) ds \geq 0$$

zusammen, wobei das Integral über die gegebene Lage der Randkurve zu erstrecken ist und an jedem Elemente ds unter w die Richtung, unter Ldw die Größe der dem Zuwachs dw entsprechenden Verrückung des Randpunktes, unter j_{AB}, j_{AC}, j_{BC} die auf ds ins Innere der Flächen hin errichteten Normalen zu verstehen sind. Da die Funktion L hier beliebig gewählt werden kann, nur daß sie stetig und stets ≥ 0 ist und an den Endpunkten der Randkurve verschwindet, so folgt hieraus

$$(7) \quad -T_{AB} \cos(wj_{AB}) - T_{AC} \cos(wj_{AC}) - T_{BC} \cos(wj_{BC}) \geq 0$$

längs der ganzen Randkurve, und zwar noch für beliebige Richtungen w . Nun können wir aber mit jeder Richtung w die entgegengesetzte kombinieren, und ist daher das Zeichen $\geq$ hier durch $=$ zu ersetzen. Hiernach müssen sich drei Vektoren von den Längen T_{AB}, T_{AC}, T_{BC} und den zu j_{AB}, j_{AC}, j_{BC} parallelen Richtungen zu einem geschlossenen Dreiecke aneinanderfügen*). Der Winkel $(j_{AB}j_{AC}) = \omega_A$ z. B. ist dann der von den zwei ersten Seiten in diesem Dreieck gebildete Außenwinkel und hat hiernach längs der ganzen Randkurve einen konstanten aus den drei Spannungen folgenden Wert. Dieser Winkel ω_A heißt der *Randwinkel* von A gegen B und C .

Ein erstes Erfordernis für das angenommene Gleichgewicht ist nun, daß aus den drei Längen T_{AB}, T_{AC}, T_{BC} überhaupt ein Dreieck zu bilden ist, d. h. daß von jenen drei Spannungen keine größer als die Summe der

*) Diese Bedingung ist von F. Neumann aufgestellt und zuerst in der Dissertation von Paul du Bois-Reymond (Berlin 1859) veröffentlicht worden.

beiden anderen ist. Ist jedoch etwa $T_{AB} > T_{AC} + T_{BC}$, so wird vielmehr C sich zwischen A und B ausziehen, eventuell zu einer dünnen Schicht mit zwei einander derart nahe liegenden Trennungsflächen gegen A und B , daß dadurch Umstände resultieren, die erst auf Grund der Annahme einer räumlichen Verteilung der Kapillarenergie genauer zu verfolgen sein werden.

Die Relation $T_{AB} > T_{AC} + T_{BC}$ für drei Medien dient als hinreichende Erklärung der mannigfachsten Kapillarphänomene*).

Nach den Beobachtungen von Marangoni**) ist in allen Fällen für zwei Flüssigkeiten die gegenseitige Oberflächenspannung kleiner als die Differenz ihrer Oberflächenspannungen gegen Luft***), hierbei also niemals jenes Dreieck von Spannungen realisierbar. Der Fall von Quecksilber und Wasser, den Marangoni als eine Ausnahme ansah, fügt sich dieser allgemeinen Regel†). Wenn Wasser auf Quecksilber in einem Tropfen steht, so haften der Quecksilberoberfläche fremde Bestandteile an, die ihre Spannung herabsetzen††).

Stellt C einen festen Körper vor, so kann die Trefflinie von A, B, C nur auf der Oberfläche dieses frei verschoben werden, und erhalten wir die Relation (7) in dem entsprechenden beschränkteren Umfange, nämlich, wenn C keine Diskontinuität der Tangentialebene an der Randkurve hat (Fig. 4), einmal so, daß w mit j_{AC} , das andere Mal so, daß w mit j_{BC} zusammenfällt, und wir erschließen damit

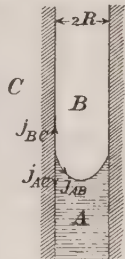


Fig. 4.

$$(8) \quad \cos \omega_A = \frac{T_{BC} - T_{AC}}{T_{AB}},$$

wo ω_A wieder den Randwinkel (j_{AB}, j_{AC}) von A gegen B und C bezeichnet†††).

*) Eine bis zur Gegenwart reichende Übersicht der Beobachtungsmethoden und -ergebnisse zur Kapillarität bringt der Artikel von F. Pockels im Handbuch der Physik, herausg. von A. Winkelmann, Bd. 1 (Breslau 1907).

**) Marangoni, Sull' espansione delle gocce di liquido galleggiante sulla superficie di altro liquido, Pavia 1865; Ann. Phys. Chem. 143 (1871), S. 348. Dieselbe Tatsache fanden van der Mensbrugghe, Mém. cour. de l'Acad. de Belg. 34 (1869); ferner Lüdtege, Ann. Phys. Chem. 137 (1869), S. 362.

***) Siehe die Anm. auf S. 297.

†) Quincke, Ann. Phys. Chem. 139 (1870), S. 66. Lord Rayleigh, Sc. papers 3, p. 562.

††) Das Ausbreiten eines Tropfens einer Flüssigkeit auf einer anderen Flüssigkeit geschieht jedesmal in charakteristischen Formen, die mit den Substanzen sehr mannigfaltig variieren, den Tomlinson'schen Kohäsionsfiguren; vgl. darüber O. Lehmann, Molekularphysik 1 (Leipzig 1888), S. 260; Paul du Bois-Reymond, Ann. Phys. Chem. 139 (1870), S. 262.

†††) Quincke (Ann. Phys. Chem. 137 (1869), S. 42) fand, daß längs einer auf Glas keilförmig aufgetragenen dünnen Silberschicht Wasser oder Quecksilber einen kon-

Diese Relation würde unmöglich sein, wenn der Quotient rechts dem Betrage nach > 1 (oder < -1) ausfällt. Im Falle $T_{BC} > T_{AC} + T_{AB}$ (es brauchten hier T_{AC} , T_{BC} nicht ≥ 0 zu sein) kommt dann Gleichgewicht dadurch zustande, daß sich die Flüssigkeit A am festen Körper C in einer selbst mikroskopisch nicht meßbaren dünnen Schicht entlang zieht, C benetzt, wodurch an der zu bemerkenden Randlinie B beiderseits an A grenzt und daher eben nach dieser Formel (8), worin nun A statt C und $T_{AA} = 0$ zu nehmen ist, sich der Randwinkel von A gleich Null herausstellt.

Hat die Wand des festen Körpers C an der Randkurve gerade eine Schneide, (ein Fall, wie er sich bei der Adhäsion einer Flüssigkeit an einem festen Körper leicht darbietet (Fig. 5)), so kommen zweierlei nicht

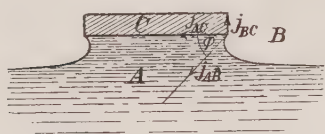


Fig. 5.

entgegengesetzte Verschiebungen der Randkurve auf C in Betracht. Das Ergebnis ist dasselbe, als wenn man sich die Schneide als Grenze abgerundeter Formen denkt; man kommt zur Ungleichung (7) einmal so, daß darin w durch j_{AC} , aber j_{BC} durch die zu j_{AC} entgegengesetzte Richtung, das andere Mal so, daß darin w durch j_{BC} , aber j_{AC} durch die zu j_{BC} entgegengesetzte Richtung vertreten wird; man erhält demnach

$$\begin{aligned} -T_{AB} \cos(j_{AC} j_{AB}) - T_{AC} + T_{BC} &\geq 0, \\ -T_{AB} \cos(j_{BC} j_{AB}) + T_{AC} - T_{BC} &\geq 0, \end{aligned}$$

d. h.

$$(9) \quad (j_{AC} j_{AB}) \geq \omega_A, \quad (j_{AB} j_{BC}) \geq \pi - \omega_A,$$

wo ω_A den durch (8) bestimmten Winkel ≥ 0 und $\leq \pi$ bedeutet. Aus beiden Relationen zusammen folgt

$$(j_{AC} j_{BC}) \geq \pi.$$

Danach kann im Zustande des Gleichgewichts die Grenzlinie der freien Oberfläche niemals längs eines endlichen Stücks auf einer konkaven Schneide des festen Körpers liegen*).

Die Bedingungen für das Zusammentreffen von vier Flüssigkeiten in einem Punkte sind nunmehr ohne weiteres ersichtlich. Die Möglichkeit der Bildung einer neuen Trennungsfläche an einer Linie, längs der mehr als drei Flüssigkeiten zusammentreffen, erörterte Gibbs**).

stanten Randwinkel erst dort ergibt, wo die Dicke der Silberlamelle mindestens 50×10^{-7} cm beträgt.

*) Gauß, Principia generalia theoriae figurae fluidorum, art. 30.

**) Gibbs, Equilibrium of heterogeneous substances, p. 453.

5. Kapillardruck. Oberflächenspannung.

Wollen wir den Begriff des Drucks in einer Flüssigkeit auch bei Erscheinungen der Kapillarität verwenden, so wird die Vorstellung notwendig, daß dieser Druck an einer Trennungsfläche zweier Flüssigkeiten im allgemeinen sich diskontinuierlich ändert. Die Diskontinuitäten sind mit den Schwerpunkts- und Flächensätzen der Mechanik in Übereinstimmung zu bringen. Will man die Diskontinuitäten weiter begründen, ohne jedoch Hypothesen über Molekularkräfte einzuführen, so kann man von dem Ansatz ausgehen, daß in einer Flüssigkeit an jeder Stelle eine räumliche Energiedichte besteht, welche von der Massendichte daselbst und auch noch von den örtlichen Differentialquotienten der Massendichte abhängt. Man hat sodann einen Grenzübergang in der Weise zu vollziehen, daß die Differentialquotienten der Massendichte im allgemeinen gleich Null gesetzt werden und nur an gewissen Flächen derart unendlich werden, daß dort die Massendichte einen konstanten Sprung erfährt. Der Begriff des Druckes entsteht dabei als der negativ genommene Differentialquotient der Energie einer Masse nach ihrem Volumen (Gl. (42) in Nr. 18).

Der Kürze wegen begnügen wir uns hier mit folgenden mehr axiomatischen Festsetzungen: Innerhalb einer einzelnen Flüssigkeit A variiert der Druck stetig mit der Dichte, ist aber nur bis auf eine additive Konstante zu bestimmen; bei gewisser Verfügung über diese Konstante wollen wir von ihm als *kinetischem Druck* P_A sprechen. Nun seien zwei verschiedene, der Schwere unterworfenen Flüssigkeiten A und B durch eine horizontale Ebene $z = 0$ getrennt und eine jede derart beschaffen, daß in ihr Dichte und Temperatur überall nur von der Vertikalhöhe z abhängen. Als dann erleidet der kinetische Druck beim Übergang von A nach B eine Diskontinuität, die wir als *Kohäsionssprung* bezeichnen wollen. Die bezügliche Abnahme des kinetischen Drucks von A nach B können wir in die Form $K_A - K_B$ setzen, so daß K_A nur von A , K_B nur von B abhängt.

Die Differenz $P_A - K_A = p_A$ soll dann der *hydrodynamische Druck* in A heißen; dieser Druck würde nun *an horizontalen Trennungsflächen keinerlei Diskontinuität* erfahren. In einer ruhenden Flüssigkeit A , in welcher die Dichte nahezu als konstant anzusehen ist, variiert der Druck p_A derart, daß er abhängig von der Vertikalhöhe z allgemein den Ausdruck $p_0 - \varrho_A g z$ hat, wo p_0 eine Konstante ist.

Nun mögen zwei ruhende Flüssigkeiten A und B von verschiedener Dichte eine beliebige, der Differentialgleichung (6) entsprechende Trennungsfläche F_{AB} haben; wir bestimmen die Niveauebene dazu, wählen sie als Ebene $z = 0$ und denken uns andererseits in einem sehr weiten Gefäße ebenfalls A und B und beide durch eine horizontale Ebene und zwar genau in

Höhe jener Niveauebene getrennt, und verbinden endlich die A beiderseits und die B beiderseits je durch eine kommunizierende Röhre (Fig. 6), so

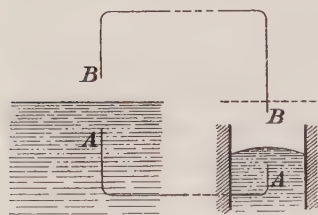


Fig. 6.

wird das Gleichgewicht nach der Gleichung (6a) bestehen bleiben. Ist nun p_0 der in jener horizontalen Trennungsfäche in A und B gleiche hydrostatische Druck, so ist dieser Druck in A in einer Höhe z gleich $p_A = p_0 - \varrho_A g z$ und in B in einer Höhe z gleich $p_B = p_0 - \varrho_B g z$. An einer beliebigen Stelle der Trennungsfäche F_{AB} findet daher gemäß (6a) eine Druckdiskontinuität

$$(10) \quad p_A - p_B = -g(\varrho_A - \varrho_B)z = T_{AB} \left(\frac{1}{R_1} + \frac{1}{R_2} \right)$$

statt. Diese Differenz heißt der *Kapillardruck* an der Stelle in A .

Stellt B den gesättigten Dampf der Flüssigkeit A vor, so ist p_0 der Sättigungsdruck über einer ebenen Flüssigkeitsoberfläche, und würde dagegen p_B den Druck des im Gleichgewicht mit der Flüssigkeit befindlichen Dampfes über einer solchen Stelle der Flüssigkeit, welche nach dem Dampf hin die mittlere Krümmung $\frac{1}{2} \left(\frac{1}{R_1} + \frac{1}{R_2} \right)$ zeigt, angeben; danach überwiegt der letztere Sättigungsdruck p_B um

$$\frac{\varrho_B}{\varrho_A - \varrho_B} T_{AB} \left(\frac{1}{R_1} + \frac{1}{R_2} \right)$$

den Druck p_0^*). Es folgt daraus z. B. ein vermehrtes Verdampfungsbestreben kleinster Wassertröpfchen in der Luft, weil mit dem Gleichgewichtsdruck des Dampfes über einer ebenen Wasseroberfläche noch nicht der Gleichgewichtsdruck über den Oberflächen der Tröpfchen erreicht ist.

Ist in dieser Weise einmal die Druckdiskontinuität des Kapillardrucks eingeführt, so würde das Bestehen der Trennungsflächenenergie nach dem Ausdrücke (3) ihrer Ableitung vollständig mit der weiteren Annahme zu erschöpfen sein, daß außerdem an jedem Randelement der Trennungsfäche in ihr, normal gegen den Rand nach innen gerichtet, eine konstante Zugspannung $= T_{AB}$, auf die Längeneinheit der Randlinie berechnet, herrscht.

Indem nun die Formel (3) sich auch auf jeden beliebigen Ausschnitt aus der Trennungsfläche anwenden läßt, würden die Kapillardrucke längs

*) W. Thomson, Edinburgh Proc. Roy. Soc. 7 (1870), p. 63. — In einer Kapillarröhre vom Radius 0,00012 cm, in welcher Wasser 1300 cm steigt, würde der Gleichgewichtsdruck des Wasserdampfes um etwa $\frac{1}{1000}$ kleiner als der Wert dafür über der Niveauebene sein. Mit den durch die Relation (10) gegebenen Umständen hängt auch der Siedeverzug luftfreier Flüssigkeiten, ferner die Schwierigkeit der Bildung der ersten Bläschen bei der Elektrolyse zusammen.

der Fläche und diese Zugspannungen an ihrem Rande für die virtuelle Arbeit gleichbedeutend mit der Annahme sein, daß *überall innerhalb der Trennungsfläche eine konstante Spannung* $= T_{AB}$ herrscht. Aber sprechen wir in solcher Allgemeinheit von einer Spannung innerhalb der ganzen Fläche, so heißt dieses im Grunde nichts anderes als: Es besteht für die Trennungsfläche eine potentielle Energie $= T_{AB} F_{AB}$, wovon wir eben ausgegangen sind.

Auf diese Analogie einer Flüssigkeitsoberfläche mit einer elastischen Haut gründete Thomas Young*) eine vollständige Theorie der Kapillarphänomene, die allerdings durch Vermeidung mathematischer Symbole an Durchsichtigkeit einbüßte. Den Begriff der Oberflächenspannung einer Flüssigkeit hat Segner**) eingeführt.

6. Formen freier Oberflächen. Tropfen.

Die Differentialgleichung einer freien Oberfläche kommt in Versuchen namentlich unter zweierlei speziellen Umständen in Betracht; nämlich es handelt sich meist entweder um Rotationsflächen um eine vertikale Achse oder aber um Zylinderflächen mit horizontalen Erzeugenden, wobei letztere Flächen auch noch als eine Approximation der ersteren bei großem Querschnitt dienen.

Im Falle einer *Rotationsfläche um die z-Achse* sei für die Meridiankurve r der Abstand von der Achse, φ der Neigungswinkel der Tangente gegen die horizontale r -Achse (Fig. 7), also $\operatorname{tg} \varphi = dz/dr$, so ist die Krümmung der Kurve $-\frac{1}{R_1} = \frac{d\varphi}{(\cos \varphi)^{-1} dr}$ und die reziproke Länge der

Normale $-\frac{1}{R_2} = \frac{\sin \varphi}{r}$, die Gleichung (6) geht also in

$$(11) \quad T_{AB} \frac{d(r \sin \varphi)}{r dr} = \lambda_{AB} + g(\varrho_A - \varrho_B)z$$

über.

Zumeist handelt es sich um eine solche partikuläre Lösung dieser Gleichung, welche die Achse trifft und sie dann notwendig senkrecht durchsetzt, damit die Fläche sich an der Stelle regulär verhält. Diese Lösung hängt nur noch von einer Konstante ab, da

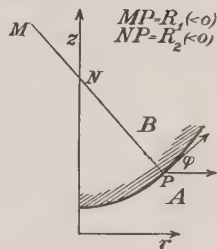


Fig. 7.

*) Th. Young, Essay on the cohesion of fluids, Phil. Trans. Roy. Soc. London 1805. — Für die Würdigung der Leistung von Young vgl. Lord Rayleigh, Phil. Mag. 30 (1890), p. 285, 456 = Scientific papers 3, p. 397.

**) Segner, Comment. soc. reg. Gotting. 1 (1751), p. 301. — Plateau (Statique des liquides, chap. V) gibt eine bis 1869 geführte historische Übersicht über die Arbeiten zur Theorie der Oberflächenspannung. Mannigfache Belege zu dieser Theorie hat namentlich Van der Mensbrugghe beigebracht.

$dz/dr = 0$ für $r = 0$ gefordert wird. Verlegt man den Koordinatenanfang in jenen Treffpunkt mit der Kurve, so bedeutet $2T_{AB}/\lambda_{AB}$ den Krümmungsradius daselbst und bei Wahl dieser Größe als Längeneinheit hängt die Form der Kurve nur noch von *einem* Parameter ab, der die Relation zwischen dem vorgeschriebenen Werte des Randwinkels von A am Ende der Meridiankurve und dem Volumen von A vermitteln muß.

Laplace*) und in der Folge Lord Kelvin**) haben die Meridiankurve der kapillaren Rotationsfläche aus kleinen Kreisbögen mit stetig sich aneinanderreihenden Tangenten unter Berechnung der Krümmung am Anfange jedes Bogens gemäß der Gleichung (11) angenähert aufgebaut. C. V. Boys***) hat diese Methode besonders handlich gemacht, indem er die Kreisbögen durch eine feste Marke an einem (durchsichtig hergestellten) Lineal beschreibt, auf dem das Drehungszentrum sukzessive verändert wird, wodurch die Stetigkeit in den Tangenten der sukzessiven Kreisbögen gesichert wird. Zudem sind die Teilstriche des Lineals durch ihre reziproken Entfernungen von der festen Marke bezeichnet, an der selbst dann ∞ steht. Bashforth†) lieferte ausgedehntes Tabellenmaterial zu jener partikulären Lösung von (11). C. Runge††) nahm die Gleichung als Beispiel bei Darlegung einer numerischen Integrationsmethode für die Differentialgleichungen zweiter Ordnung. Eine im Bereiche $0 \leq \varphi < \pi/2$ konvergente Entwicklung von z nach Potenzen von r für die Lösung von (11) behandelten K. Lasswitz†††), Th. Lohnstein*†). Allerhand Annäherungsformeln, bzw. des Krümmungsradius für $r = 0$, des Maximalwertes von r usw. findet man bei Poisson**†), Fr. Neumann***†), A. König*††), H. Siedentopf**††).

Die Formen eines Quecksilbertropfens auf einer horizontalen Unterlage, einer gegen eine Horizontalebene stoßenden Luftblase, eines an einer Horizontalebene hängenden Wassertropfens sind Rotationsflächen, bestimmt

*) Laplace, *Connaissance des Temps*, 1812.

**) W. Thomson, *Capillary attraction*, Proc. Roy. Inst. 11 (1886), aufgen. in *Popular lectures and addresses* 1, London 1889. Der Aufsatz enthält verschiedene Diagramme zur Illustrierung des Verfahrens. — J. C. Schalkwijk, *Leiden Communic.* No. 67 (1901).

***) C. V. Boys, *Phil. Mag.* (5) 36 (1893), p. 75.

†) Bashforth and Adams, *An attempt to test the theories of capillary action*, Cambridge 1883.

††) C. Runge, *Math. Ann.* 46 (1895), S. 167.

†††) K. Lasswitz, *Inaug.-Diss.* Breslau 1873.

*†) Th. Lohnstein, *Inaug.-Diss.* Berlin 1891.

**†) Poisson, *Nouv. théor. de l'act. capill.*, Paris 1831.

***†) Fr. Neumann, *Vorl. über Capill.* 1894.

*††) A. König, *Ann. Phys. Chem.* 16 (1882), S. 10.

**††) H. Siedentopf, *Ann. Phys. Chem.* 61 (1897), S. 235.

durch die Differentialgleichung (11), durch die Forderung, die Achse zu treffen, durch den Randwinkel am Endpunkt der Meridiankurve und durch das vorliegende Volumen.

Hängt die Lösung der Gleichung (6) von y nicht ab, ist sie also eine *Zylinderfläche mit horizontalen Erzeugenden parallel der y -Achse*, so wird die Gleichung ihres vertikalen Querschnitts mit der xz -Ebene, wenn φ den Winkel der Tangente gegen die x -Achse, ds das Bogenelement bedeutet:

$$(12) \quad T_{AB} \frac{d \sin \varphi}{dx} = T_{AB} \frac{d \varphi}{ds} = \lambda_{AB} + g(\varrho_A - \varrho_B)z.$$

Es ist das die Gleichung der Gleichgewichtsform, die ein elastischer gleichförmiger unendlich dünner und ohne äußere Kräfte geradliniger Stab annimmt, wenn an den Enden zwei in die Richtung der positiven und negativen x -Achse fallende entgegengesetzt gleiche Kräfte und dazu die geeigneten Kräftepaare angreifen*). Die Differentiation von (12) nach s ergibt

$$T_{AB} \frac{d^2 \varphi}{ds^2} = g(\varrho_A - \varrho_B) \sin \varphi,$$

und ist danach, $\varrho_A > \varrho_B$ angenommen, die Abhängigkeit des Winkels $\pi - \varphi$ von s dieselbe wie des Ausschlags eines gewöhnlichen mathematischen

Pendels mit der Länge $\frac{T_{AB}}{\varrho_A - \varrho_B}$ von der Zeit. Wird $z = 0$ nach der Niveauebene gelegt, also $\lambda_{AB} = 0$ angenommen, so erhält man aus (12) durch Multiplikation mit $\operatorname{tg} \varphi dx = dz$ und Integration, entsprechend dem Integral der lebendigen Kraft in der Pendelbewegung:

$$(13) \quad T_{AB}(c - \cos \varphi) = g(\varrho_A - \varrho_B) \frac{z^2}{2}.$$

Darin ist die Integrationskonstante $c = 1$, wenn die Fläche sich asymptotisch an die Niveauebene heranzieht, und $c > 1$, wenn sie sonst eine horizontale Tangente hat; andererseits ist, wenn die Fläche einen Wendepunkt ($\frac{d\varphi}{ds} = 0$) besitzt, notwendig $c < 1$.

Die Form eines an einer Horizontalebene hängenden zylindrischen Tropfens, wie er durch Austreten einer Flüssigkeit aus einem langen Spalt entstehen könnte, hat Fr. Neumann**) behandelt. Um über die Stabilität der Form zu entscheiden, haben wir das Jacobische Kriterium für ein Extremum in einem Variationsproblem heranzuholen. Benetzt A die Ebene und wird der Einfachheit halber $2T_{AB}:g(\varrho_A - \varrho_B)$ als Flächeneinheit eingeführt, so kommt hier das Variationsproblem darauf hinaus,

*) Vgl. z. B. A. E. H. Love, A treatise on the mathematical theory of elasticity 2 (Cambridge 1893), Arts. 227—229. [(2nd edition 1906, Art. 262. Deutsche Ausgabe, bes. v. A. Timpe (Leipzig 1907) §§ 262, 263)].

**) Fr. Neumann, Vorl. über Capill., S. 117.

in einem Intervalle $-x_0 \leq x \leq x_0$, dessen Länge $2x_0$ ebenfalls noch gesucht wird, eine an den Enden verschwindende stetige Funktion $z(x)$ derart zu bestimmen, daß

$$\int_{-x_0}^{x_0} \left(\sqrt{1 + \left(\frac{dz}{dx} \right)^2} - 1 - z^2 \right) dx$$

zu einem *Minimum* wird, während $\int_{-x_0}^{x_0} z dx = J$ gegeben ist. In einer ge-

wissen Tiefe $\frac{1}{2}z_0$ unter der Horizontalebene zeigt das tellerförmige Profil des Tropfens durch einen Wendepunkt die Niveauebene an und verläuft sodann gemäß (13) bis zum tiefsten Punkte als spiegelbildliche Fortsetzung am Wendepunkt, so daß z_0 die ganze Tiefe des Tropfens wird. Ist 2θ die Neigung der Wendetangente gegen die Horizontale und $\kappa = \sin \theta$, so findet man auf Grund von (13):

$$z_0 = 2\sqrt{2}\kappa, \quad x_0 = \sqrt{2}(2E - K), \quad J = x_0 z_0,$$

wo K und E die vollständigen elliptischen Integrale erster und zweiter Gattung vom Modul κ sind. Der Ausdruck $J = x_0 z_0$ hat ein Maximum ungefähr bei $\theta = 35^\circ 32'$ mit $J = 2,606$. Nur wenn das Volumen des Tropfens auf die Längeneinheit des Spalts, J , unterhalb dieser Größe liegt, gibt es überhaupt Tropfenformen, welche den Gleichungen des Problems entsprechen, und zwar dann eine breitere und weniger tiefe Form, wobei $\theta < 35^\circ 32'$ ist, und eine schmalere tiefer herunterhängende, für welche diese stärkste Neigung gegen die Horizontale $> 35^\circ 32'$ ist. Nur die erste Form ist stabil.

Daß für rotationsförmige hängende Tropfen die Verhältnisse analog liegen dürften, geht aus einem Experiment von Lord Kelvin*) hervor, wonach eine um einen horizontalen Metallring gespannte dünne Kautschukhaut, die durch Hinaufgießen von Wasser in eine tropfenähnliche Form gedehnt wird, in einem gewissen Stadium der Füllung ruckweise eine Lage instabilen Gleichgewichts passiert.

Nach dem Abreißen eines Tropfens zieht sich der ausgezogene zurück-schnellende Hals in einen oder mehrere kleinere Tropfen zusammen. Der Vorgang wird der Beobachtung zugänglicher, wenn die Tropfenbildung in einer nur wenig leichteren Flüssigkeit erfolgt, ist jedoch einer mathematischen Behandlung noch nicht unterzogen**).

*) W. Thomson, Popular lectures and addresses 1, London 1889, p. 38. — Daraus sind die Tropfenformen in Fig. 8 oben entnommen.

**) G. Hagen, Ann. Phys. Chem. 67 (1846), S. 1, 152; 77 (1849), S. 449. — C. V. Boys, Seifenblasen. Vorl. über Capill. Deutsche Übers. von G. Meyer, Leipzig 1893, S. 33, 65. — Die Beziehungen zwischen dem Durchmesser einer Röhre und dem Gewicht daraus abfallender Tropfen behandeln Lord Rayleigh, Phil. Mag. 48 (1899),

7. Steighöhen.

Die in einem Gefäße C senkrecht unterhalb der Trennungsfläche F_{AB} , von der Niveaubene $z = z_{AB}$ an gerechnet, stehende Masse von A überwiegt die dadurch verdrängte Masse von B um

$$(14) \quad \begin{aligned} g(\varrho_A - \varrho_B) \int_{F_{AB}} (z - z_{AB}) \cos(nz) df &= - T_{AB} \int_{F_{AB}} \left(\frac{1}{R_1} + \frac{1}{R_2} \right) \cos(nz) df \\ &= - T_{AB} \int \cos(j_{AB}z) ds, \end{aligned}$$

wo letzteres Integral sich über den Rand von F_{AB} erstreckt. Die erste Umformung folgt aus (6), die zweite durch Anwendung der Formel (3) auf eine Parallelverschiebung der Fläche in der z -Richtung, wobei ihr Flächeninhalt sich nicht ändert. Steht die Gefäßwand am Rande von F_{AB} überall vertikal, so ist hier $(j_{AB}z) = \pi - \omega_A$, unter ω_A den Randwinkel von A verstanden, und wird daher der letzte Ausdruck in (14) $= T_{AB} \cos \omega_A U$, wo U den Umfang der Randkurve bedeutet, insbesondere demnach positiv, Null oder negativ, je nachdem der Winkel ω_A spitz, ein rechter oder stumpf ist.

Stellt C eine vertikale Kapillarröhre mit kreisförmigem Querschnitte vom Radius R vor, so tritt in der Röhre ein Aufsteigen oder eine Depression von Flüssigkeit ein (wir denken uns hier $\varrho_A > \varrho_B$ und B oberhalb A gelegen), je nachdem der Randwinkel von A spitz oder stumpf ist, im speziellen also ein Ansteigen, wenn A die Röhre benetzt. Die *mittlere Steighöhe* über der Querschnittsfläche der Röhre ist nach (14)

$$h_m = \frac{2 T_{AB} \cos \omega_A}{g(\varrho_A - \varrho_B) R} = 2 \frac{T_{BC} - T_{AC}}{g(\varrho_A - \varrho_B) R},$$

also *umgekehrt proportional dem Radius der Röhre**). Der Meniskus läßt sich in erster Annäherung als eine Kugelfläche ansehen. Approximiert man ihn genauer als ein Rotationsellipsoid um die Röhrenachse**), welches mit ihm im Randwinkel, in dem Krümmungsradius auf der Achse und im

p. 321 (Sc. papers 4, p. 415), Th. Lohnstein, Ann. Phys. Chem. 20 (1906), S. 237, S. 606. — A. M. Worthington and R. S. Cole, Impact with a liquid surface, London Phil. Trans. 189 (1897), p. 187.

*) Die Proportionalität der Steighöhe in einer Kapillarröhre mit dem Reziproken des Durchmessers scheint zuerst von Borelli (De motionibus naturalibus a gravitate pendentibus, Reggio 1670) ausgeführt zu sein; der Satz wird von manchen Autoren Jurin (Phil. Trans. 30 (1718)) zugeschrieben.

**) Mathieu, Capillarité, Paris 1883, p. 49.

angehobenen Gewicht übereinstimmt, so folgt z. B., wenn A die Röhre benetzt, als Steighöhe auf der Achse

$$h = \frac{h_m}{1 + \frac{1}{3} \frac{R}{h_m}}.$$

Werden in einer Kapillarröhre mehrere die Wand nicht benetzende Flüssigkeiten $A, B, B^*, \dots$ übereinander geschichtet, so ist das gesamte angehobene Gewicht das nämliche, als wenn sich über A nur B befände. Ein Einwand, den Young aus Beobachtungen gegen diese Schlußfolgerung und damit überhaupt gegen die Theorie von Laplace erheben zu müssen glaubte, wurde durch Poisson*) entkräftet.

Zwischen zwei parallelen vertikalen Platten ist zufolge (14) die mittlere Steighöhe halb so groß als in einer Kapillarröhre von einem Durchmesser gleich dem Abstand der Platten. Stehen die zwei vertikalen Platten mit geringer Keilöffnung gegeneinander, so steigt die Flüssigkeit an ihnen zu

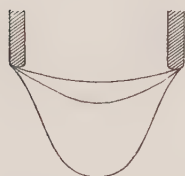


Fig. 8.

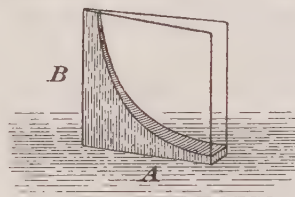


Fig. 9.

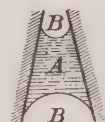


Fig. 10.



Fig. 11.

einer gleichseitigen Hyperbel empor (Fig. 9). In einer konischen Röhre kann unter Umständen ein Tropfen im Gleichgewicht sein bei spitzem Randwinkel, wenn die Röhre sich nach oben verjüngt (Fig. 10), oder bei stumpfem Randwinkel, wenn sie sich nach unten verjüngt (Fig. 11).

8. Kapillarauftrieb. Adhäsion.

Der Körper C sei nur mit den Flüssigkeiten A und B in Berührung. Um den von C zur Erhaltung des Gleichgewichts gegen A und B zu leistenden Gegendruck in der Komponente $-P_w$ nach einer beliebigen Richtung w zu ermitteln, lassen wir C in dieser Richtung parallel mit sich verschiebbar sein. Wir nehmen sodann eine von einem Parameter w abhängende Schar von Verrückungen des Systems vor, wobei C in jener Richtung um die Längen w fortschreitet, die Partien F_{AC} , F_{BC} , also auch ihre gemeinsame Randlinie unverändert mitgehen, alle Grenzflächen in denen A und B an andere Medien als C anstoßen, festbleiben, endlich F_{AB} noch sich derart deformiert, daß die Volumina V_A und V_B ungeändert bleiben.

*) Poisson, Nouv. théor. de l'act. capill., p. 141.

Wir können alsdann für die gesamte ins Spiel kommende Energie E , einschließlich des Terms wP_w für den Gegendruck $-P_w$, die Relation $\frac{dE}{dw} = 0$ ansetzen. Nun sind die Flächeninhalte von F_{AC} , F_{BC} unverändert, längs F_{AB} besteht die Differentialgleichung (6), zur Vereinfachung legen wir $z = 0$ in die Niveaubene von A , B , haben also $\lambda_{AB} = 0$. Im Hinblick auf (3) und (5) erhalten wir daher:

$$(15) \quad P_w - \int_{F_{AC}} g(\varrho_A - \varrho_C) z \cos(w n) df - \int_{F_{BC}} g(\varrho_B - \varrho_C) z \cos(w n) df \\ - T_{AB} \int \cos(w j_{AB}) ds = 0,$$

wo sich das erste Integral auf F_{AC} , das zweite auf F_{BC} , das dritte auf ihren gemeinsamen Rand bezieht und n die äußere Normale von C bezeichnet.

Der auf C ausgeübte vertikale Auftrieb berechnet sich hieraus, indem wir für w die z -Richtung nehmen. Hat die Trennungsfläche F_{AB} keine Begrenzung außer ihrer Randlinie auf C , d. h. verläuft sie im übrigen asymptotisch an die Niveaubene, so zeigt die bei (14) vorgenommene Transformation, daß der letzte Term in (15) alsdann $= g(\varrho_A - \varrho_B) V_{AB}$ wird, unter V_{AB} das unterhalb F_{AB} bis zur Niveaubene reichende Volumen verstanden (soweit F_{AB} unterhalb der Niveaubene verläuft, ist das dazwischenliegende Volumen in V_{AB} negativ einzurechnen). Von diesem Volumen entfalle der Anteil V auf das Medium A (Fig. 12). Andererseits werde C durch Fortführung der Niveaubene in einen unteren Teil vom Volumen $V_C^{(A)}$ und einen oberen Teil vom Volumen $V_C^{(B)}$ zerlegt, so werden der zweite und dritte Term in (16) bzw.

$$-g(\varrho_A - \varrho_C)(V_C^{(A)} + V_{AB} - V), \quad -g(\varrho_B - \varrho_C)(V_C^{(B)} - V_{AB} + V)$$

und folgt demnach

$$(16) \quad P_z = g(\varrho_A - \varrho_C) V_C^{(A)} + g(\varrho_B - \varrho_C) V_C^{(B)} - g(\varrho_A - \varrho_B) V.$$

Die ersten zwei Terme bilden den hydrostatischen Auftrieb, falls die Trennungsfläche in die Niveaubene fiele, der dritte Term, der *kapillare Auftrieb* (bzw. negative Abtrieb) ist entgegengesetzt gleich dem infolge der Kapillarität über die Niveaubene gehobenen Flüssigkeitsgewicht. Hiernach kann bei *stumpfen* Randwinkel ω_A unter Umständen ein Körper auf einer Flüssigkeit von geringerem spezifischen Gewicht schwimmen.

Wird eine kreisförmige Scheibe C auf eine weite horizontale Oberfläche von A in B gelegt (Fig. 5) und mit horizontal bleibender Basis,

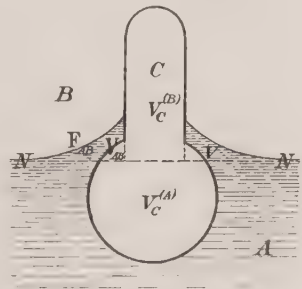


Fig. 12.

die stets ganz mit A in Berührung sei, kontinuierlich senkrecht gehoben, so entspricht die am Rande der Scheibe ansetzende freie Rotationsfläche wieder der Gleichung (11); die Meridiankurve verläuft asymptotisch an die Niveaubene, während der Randwinkel φ von A gegen die horizontale Basis der Scheibe kontinuierlich abnehmend zufolge der ersten Ungleichung (9) nur bis zu dem durch (8) bestimmten Werte ω_A heruntergehen kann, wobei dann die Flüssigkeit abreißt. Bei großem Flächeninhalt S der Scheibe ergibt sich die maximale Höhe z_0 des Anhebens angenähert aus (13) für $c = 1$, $\varphi = \omega_A$ und zwar als unabhängig von S und folgt das dabei über die Niveaubene gehobene maximale Flüssigkeitsgewicht + dem Gewicht der Scheibe aus (16), indem dort $V_c^{(A)} = -z_0 S$ substituiert und $V = V_{AB}$ aus (14) mittels $(j_{AB}z) = \frac{\pi}{2} + \omega_A$ berechnet wird.

Bei der Adhäsion zweier sehr nahe befindlicher gleicher horizontaler Platten, sie mögen etwa wieder kreisförmig vom Flächeninhalte S sein, vermöge einer zwischen ihnen befindlichen dünnen und sie *benetzenden* Flüssigkeitsschicht A vom Volumen V_A ist für die angenähert durch (12) bestimmte Meridiankurve die Höhe z , die wir von der oberen Fläche der Schicht rechnen, und damit auch $d\varphi/ds$ wenig veränderlich und daher die Kurve angenähert ein Halbkreis vom Durchmesser V_A/S (Fig. 13). Nach (12) befindet sich dann die Niveaubene in einer Höhe $z = -z_0$, die umgekehrt proportional diesem Werte ist. Aus (15) entsteht wieder genau die Relation (16), wobei $V_c^{(A)} = -Sz_0$, $V = 0$ einzusetzen ist, und folgt daraus der auf die obere Platte mitsamt ihrem Gewicht ausgeübte Zug nach unten, und zwar als proportional mit S^2/V_A . Bei kleinem V_A kann daher eine äußerst große Kraft zur Trennung der Platten nötig sein.

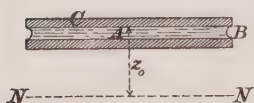


Fig. 13.



Fig. 14.

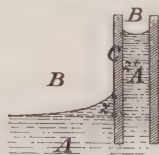


Fig. 15.

Sind dagegen die *Randwinkel* an den Platten *stumpf*, so liegt die Niveaubene über der Schicht, es richtet sich die Größe der von A bedeckten Fläche der Platten nach dem Werte von V_A , und es ist ein entsprechender Druck auf die Platten nötig, um ihre Distanz zu verringern (Fig. 14).

Stellt C eine auf beiden Seiten gleich beschaffene vertikale, der yz -Ebene parallele Platte von einer sehr großen Breite L vor, die in A und B eintaucht, wobei aber der Stand der Trennungsflächen F_{AB} beiderseits an C verschieden hoch sein kann (Fig. 15), so berechnet sich aus (15) die

Summe der zwei Drucke P_x^- und P_x^+ in Richtung der x -Achse, welche die Platte links und rechts, auf der Seite der kleineren bzw. der größeren x erfährt; die zwei Randintegrale heben sich auf, die Flächenintegrale bleiben nur für den einerseits von A , andererseits von B bedeckten Teil der Platte übrig. Steht A links bis zur Höhe z^- , rechts bis zur Höhe z^+ an der Platte, so resultiert als Gesamtdruck

$$P_x = g(\varrho_A - \varrho_B) \frac{(z^+)^2 - (z^-)^2}{2} L = T_{AB}(c^+ - c^-) L,$$

wenn für die Form der Fläche F_{AB} nach (13) links die Integrationskonstante c^- , rechts c^+ in Betracht kommt.

Tauchen jetzt zwei Platten C^- und C^+ von gleicher Breite L parallel zur yz -Ebene und sehr nahe zueinander ein und kommt für den Meniskus in der xz -Ebene zwischen ihnen die Integrationskonstante c in Betracht, während jenseits von ihnen die Flächen F_{AB} asymptotisch an die Niveauebene verlaufen mögen, also hier die betreffende Konstante den Wert 1 hat, so werden die Platten mit einer Kraft $T_{AB}(c - 1)L$ gegeneinander getrieben. Bildet nun A an beiden Platten spitze oder an beiden Platten stumpfe Winkel, so zeigt der Meniskus in der xz -Ebene zwischen den Platten notwendig eine Stelle mit horizontaler Tangente und ist daher nach (13): $c > 1$, es findet also eine scheinbare Anziehung der Platten statt und zwar, da $c - 1$ nach (13) dem Quadrat der Steighöhe an jener Stelle proportional ist, angenähert umgekehrt proportional dem Quadrat des Abstandes der Platten. Bildet A an einer Platte spitze, an der anderen stumpfe Winkel, so entspricht einer gewissen Distanz der Platten ein labiles Gleichgewicht, das bei Annäherung der Platten (durch stärkere Krümmung des Meniskus) zu einer Anziehung, bei Entfernung zu einer Abstoßung führt*).

9. Ausschaltung der Schwerkraft.

Die Wirkung der Schwere auf die Gestalt der Trennungsfläche von A gegen B erscheint nach (6) ausgeschaltet, wenn $\varrho_A = \varrho_B$ ist, die beiden Flüssigkeiten also gleiche Dichte haben. Dieser Umstand, an den schon Segner**) gedacht hat, wurde von Plateau***) vielfältig benutzt, um reine Kapillarkwirkungen zu studieren.

*) Laplace, Suppl. à la théor. de l'act. capill. (De l'attraction et de la repulsion apparente des petits corps qui nagent à la surface des fluides). — Poisson, Nouv. théor. de l'act. capill., chap. VI. — Allgemeinere Theoreme über Anziehung und Abstoßung schwimmender Körper entwickelt W. Voigt, Compendium der theor. Phys. 1, Leipzig 1895, S. 239.

**) Siehe Anm. S. 311.

***) Plateau, Mém. de l'Acad. de Belgique, 1843 bis 1868; Statique expérimentale et théorique des liquides (Gand 1873).

Ein Öltropfen, in eine gleich schwere Mischung von Wasser und Alkohol gebracht, nimmt nach (6) im Gleichgewicht die Figur einer Fläche konstanter mittlerer Krümmung an. Schwebt der Tropfen vollkommen frei, so zeigt er daher notwendig Kugelgestalt, denn die Kugel ist die

einzige geschlossene singularitätenfreie Fläche von konstanter mittlerer Krümmung*). Ist die Oberfläche des Tropfens nicht allseitig geschlossen, sondern lehnt sie sich teilweise an eingetauchte Rotationskörper an, so mag sie sich als eine Rotationsfläche um die bezügliche Achse bilden. Sind nun auf einer beliebigen Normale der Meridiankurve dieser Fläche nacheinander (Fig. 16) P der Punkt der Kurve, M das Krümmungszentrum, N der Treffpunkt mit der Achse, also

PM , PN die zwei Hauptkrümmungsradien der Rotationsfläche und ist endlich Q derart gelegen, daß $PNQM$ vier harmonische Punkte sind, also

$$\frac{1}{PM} + \frac{1}{PN} = \frac{2}{PQ}$$

ist, so muß nach (6) oder (11) die Länge PQ konstant ausfallen; das Spiegelbild Q^* von Q an der Achse liefert daher eine konstante Summe $PN + NQ^* = PQ$, während PN und NQ^* entgegengesetzt gleiche Neigung gegen die Achse zeigen. Lassen wir nun P die Meridiankurve beschreiben und konstruieren fortwährend in der dargelegten Weise N , Q , Q^* , so wird, weil PQ konstant ist, Q eine Parallelkurve zur Meridiankurve beschreiben, daher die Bewegung von Q stets normal zu NP und also die spiegelbildlich dazu an der Achse verlaufende Bewegung von Q^* stets normal auf NQ^* sein. Daraus ist ersichtlich, daß die Meridiankurve unserer Rotationsfläche durch den Brennpunkt P eines bestimmten Kegelschnittes erzeugt wird, den man ohne Gleiten auf der Rotationsachse abrollen läßt, dessen anderer Brennpunkt Q^* und dessen doppelte große Achse PQ ist**).

Wird der Tropfen durch zwei mit den Zentren vertikal übereinander liegende horizontale Scheiben oder Ringe gestützt, so können durch Abänderung der Distanz dieser Stützen sowie der zwischen ihnen befindlichen Ölmasse die verschiedenen Formen dieser Rotationsflächen konstanter mittlerer Krümmung erzielt werden, das *Unduloid* in den Grenzen Kugel — Zylinder — Katenoid, das *Nodoid* in den Grenzen Katenoid — Kugel, welche bzw. einer rollenden Ellipse oder Hyperbel und den Grenzflächen Strecke,

*) Vgl. Liebmann, Math. Ann. 53, S. 81.

**) Ch. Delaunay, J. de math. (1) 6 (1841), p. 309. — Die Literatur über die Flächen mittlerer Krümmung bis 1869 bespricht ausführlich Plateau (*Statique des liquides* 1, p. 131).

Kreis, Parabel entsprechen (Fig. 17). Dabei werden außerhalb an den Ringen sich jedesmal noch Kugelkalotten von der nämlichen mittleren Krümmung wie der dazwischen befindliche Tropfen ansetzen. Das Katenoid ist hier eine stabile Gleichgewichtsfigur, nämlich wirklich eine Fläche von kleinstem Flächeninhalt bei gegebener Größe des zwischen den zwei Basiskreisflächen gehaltenen Volumens, nur so lange die Tangenten in den zwei Endpunkten der sie erzeugenden Meridiankurven ihren Schnittpunkt vor der Rotationsachse finden*), und die Zylinderfläche ist in demselben Sinne stabil, nur so lange die Höhe des Zylinders nicht den Umfang des Querschnitts erreicht**).

Wird der freischwebende und eine Kugelgestalt bildende Öltropfen mit Hilfe einer in das Öl eingetauchten Scheibe in gleichförmige Rotation um eine Achse — etwa die z -Achse — gesetzt, so entsprechen wachsenden Werten der Winkelgeschwindigkeit ω der Rotation verschiedene Gestalten des Tropfens; er erscheint zuerst ellipsoidisch, vertieft sich oben und unten, endlich löst sich am Äquator ein Ring ab, der an der Rotation teilnimmt***). Denkt man sich, was freilich dem Versuche nur unzureichend entspricht, es rotiere nur der Tropfen A , nicht die umgebende Flüssigkeit B , und behandelt die Bewegung von mitrotierenden Koordinatenachsen aus unter Einführung des Potentials der Zentrifugalkräfte $-\frac{\omega^2}{2} \varrho_A \int_A r^2 dv$, so erhält man (mit Bezeichnungen wie in (11)) als Gleichung für die Meridiankurve des rotierenden Tropfens:

$$T_{AB} \frac{d(r \sin \varphi)}{r dr} = \lambda_{AB} - \frac{\omega^2}{2} \varrho_A r^2 \quad \left(\frac{dz}{dr} = \operatorname{tg} \varphi \right).$$

Hieraus bestimmt sich z als hyperelliptisches Integral in r vom Geschlecht 2 und kommt man je nach den Werten von ω auf sphäroidische oder ringförmige Flächen†). — Einen Vergleich der hier auftretenden Figuren mit den Gestalten gravitierender in stationärer Rotation befindlicher Flüssigkeitsmassen könnte man allenfalls zustande bringen, indem man von Fernkräften mit dem Ausdrücke $-k \frac{e^{-cr}}{r}$ ($c \geq 0$) als Potential für zwei Masseneinheiten in der Distanz r ausgeht, woraus einerseits

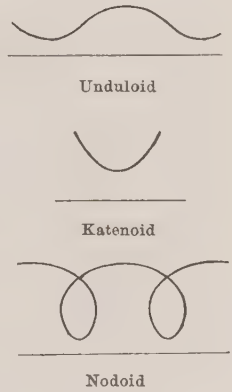


Fig. 17.

*) L. Lindelöf in Moigno-Lindelöf, Calcul des variations, Paris 1861, p. 209, 231. — Poincaré, Capillarité, p. 66.

**) Plateau, Statique des liquides 2, chap. IX. — Poincaré, Capillarité, p. 95.

***) Plateau, Mém. de l'Acad. de Bruxelles 16 (1843).

†) Beer, Einl. in die math. Theorie der Elastizität u. Kapillarität, Leipzig 1869.

Gravitation, andererseits Oberflächenspannung als die zwei Grenzfälle $c = 0$ und $c = \infty$ folgen.

10. Flüssigkeitshäute.

Unter Umständen kann eine Flüssigkeit A in einem Medium B längere Zeit hindurch als eine dünne Haut mit zwei einander sehr nahen Trennungsflächen gegen B bestehen. Die Dauerhaftigkeit solcher Flüssigkeitshäute beruht nach Plateau*) auf einer vornehmlich nach den Grenzschichten hin hervortretenden gallertartigen Beschaffenheit (Oberflächenviskosität). Diese wieder erklärt sich durch eine andere Verteilung der stofflichen Bestandteile in den Oberflächenschichten als im Inneren der Haut, wodurch jene Schichten eher die Eigenschaften eines festen Körpers als einer Flüssigkeit haben**). In Häuten von sehr geringer Dicke wird dann ein Fließen des Inneren zwischen den Oberflächenschichten außerordentlich durch die innere Reibung der Flüssigkeit verzögert***)) und dadurch eine Variation des gegenseitigen Abstandes der zwei Trennungsflächen sehr erschwert. Sind F_{AB}^- , F_{AB}^+ die Flächeninhalte der zwei Seiten der Haut, so ist alsdann zum Gleichgewicht der Haut das Minimum der potentiellen Energie

$$T_{AB}(F_{AB}^- + F_{AB}^+) + g(\varrho_A - \varrho_B) \int_A z dv + g\varrho_B \int_{A+B} z dv$$

ganz allein in bezug auf solche virtuelle Verrückungen von A zu fordern, wobei die normalen Abstände der zwei Trennungsflächen ungeändert bleiben; denn andere Verrückungen sind als unausführbar anzusehen. Diese Forderung kommt nun, wenn noch die Dicke der Haut als verschwindend zu betrachten ist, im Hinblick auf (3), (4), (5) einfach darauf hinaus, daß für die Haut $2T_{AB}F_{AB}$ oder also F_{AB} , darunter die ganze Ausdehnung der Haut verstanden, ein Minimum sein soll. Wird z. B. ein Rahmen, irgendwie aufgebaut aus festen Drähten, beweglichen Fäden, als Stütze dienenden festen Oberflächen, in eine Seifenlösung getaucht, so spannt sich hiernach innerhalb der vorgeschriebenen festen und veränderlichen Grenzen die Seifenlösung in der Form einer *Minimalfläche* (Artikel von Lilienthal, Enzyklopädie III D 5, S. 307) aus, und man trifft hier einen der seltenen Fälle an, daß ein rein mathematisches Gebiet aus einer verhältnismäßig leichten Experimentierkunst die vielseitigste Anregung zu schöpfen vermocht hat.

*) Plateau, Statique des liquides 2, chap. VII.

**) Marangoni, Nuovo Cimento (2) 5, 6 (1871/72); (3) 3 (1878). — Lord Rayleigh, Proc. Roy. Soc. 48 (1890), p. 127 (Sc. papers 3, p. 363).

***)) Vgl. die bezüglichlichen Rechnungen bei Gibbs, Equilibrium of heterogeneous substances, p. 475.

Als Differentialgleichung für die Form der Haut erhält man

$$(17) \quad \frac{1}{R_1} + \frac{1}{R_2} = 0,$$

während für ihren Rand, soweit er nicht fest vorgeschrieben ist, die Bedingung resultiert, auf die dazu dargebotenen Flächen senkrecht aufzutreffen. Dabei können sich infolge der vorgeschriebenen Grenzbedingungen Kreuzungsstellen der Haut in ihrem Verlaufe als notwendig erweisen; in stabilem Gleichgewicht können aber niemals mehr als drei Lamellen längs einer Kurve und zwar dann immer nur unter gleichen Flächenwinkeln, also von 120° , zusammentreffen und höchstens vier, und zwar mit gleichen Raumwinkeln um einen Punkt herum ansetzen*). So bildet sich z. B. in dem Kantengerüst eines regulären Tetraeders eine Seifenhaut, bestehend aus sechs ebenen Lamellen, den sechs Dreiecken vom Schwerpunkt des Tetraeders aus nach den einzelnen Kanten, dagegen entsteht innerhalb des Kantengerüsts eines Würfels jedesmal eine Fläche, die nicht alle Symmetrien des Würfels übernimmt, sondern ein beliebiges Paar seiner Seitenflächen begünstigt (Fig. 18**).

Es ist eine charakteristische Eigenschaft der Minimalflächen, daß ihre Abbildung durch parallele Normalen auf eine Kugelfläche eine konforme mit Umlegung der Winkel ist. Soll nun die Begrenzung der Minimalfläche ein gegebener geschlossener Streckenzug sein oder allgemeiner, soll sie stückweise in vorgeschriebenen Geraden oder Ebenen verlaufen, so müssen die Geraden Asymptotenkurven auf der Fläche werden und die Ebenen Krümmungskurven aus ihr heraus schneiden, und jene sphärische Abbildung wird ein Kreisbogenpolygon von bekanntem Umriß. Die analytische Bestimmung der fraglichen Minimalfläche erfordert die konforme Abbildung dieses Polygons auf eine Halbebene, diese Abbildungsaufgabe hängt von einer linearen Differentialgleichung zweiter Ordnung mit rationalen Funktionen als Koeffizienten ab, und schließlich soll man eine endliche Anzahl von Parametern, die in diese Gleichung eingehen, den Längen und Winkeln des gegebenen Rahmens entsprechend einrichten, worin transzendente Relationen liegen, deren Theorie erst in Spezialfällen zu befriedigendem Abschluß gebracht werden konnte***).

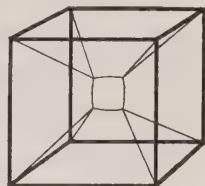


Fig. 18.

Plateau†) und in der Folge H. A. Schwarz***) haben eine Menge verschiedenartiger Minimalflächen (z. B. das Katenoid innerhalb zweier senk-

*) Lamarle, Mém. de l'Acad. de Belg. 35, 36. — Plateau, Statique des liquides 1, chap. V. **) Plateau, l. c. p. 318.

***) Vgl. H. A. Schwarz, Gesammelte math. Abh. 1, Berlin 1890.

†) Siehe Anm. S. 319.

recht übereinander gehaltener Kreisringe, eine Schraubenfläche innerhalb eines Glaszylinders zwischen zwei Erzeugenden) durch Seifenlamellen realisiert und zugleich die Grenzen ihres extremalen Charakters sowie die Umlagerungen bei eintretender Instabilität theoretisch wie experimentell festgestellt. — Innerhalb eines Drahtes, der in den sechs Kanten eines geraden, regelmäßigen, sechsseitigen Prismas und den sie abwechselnd in der einen und der anderen Grundfläche verbindenden Seiten ausgespannt ist, bildet sich, falls die Prismenkanten im Verhältnis zu den Basisseiten hinreichend lang sind, eine Lamelle aus, die auf der Mittellinie des Prismas einer der zwei Grundflächen wesentlich näher liegt und die durch ein leichtes Schütteln in das Gegenbild in bezug auf die andere Grundfläche überspringt; diese auffallende Erscheinung soll aber ganz allein auf die stets vorhandenen geringen Unvollkommenheiten der Modelle zu schieben sein.

Für die Stabilität einer Flüssigkeitshaut in einem festen Rahmen ist die Bedingung die, daß keine unendlich nahe Minimalfläche durch irgendein auf der Fläche liegendes geschlossenes Kurvensystem möglich ist*). Für den Fall beweglicher Grenzen geben die allgemeinen Kriterien von Hilbert**) bezüglich des Vorhandenseins eines Extremums Aufschluß über die Stabilität.

Indem man geeignet verfährt, kann man innerhalb eines festen Rahmens auch Seifenlamellen einspannen, in denen vollständig geschlossene Flächen (Blasen) auftreten. (Zum Beispiel kann man innerhalb des Kantengerüsts eines Würfels eine Seifenhaut herstellen, welche aus einer innen schwebenden

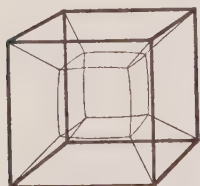


Fig. 19.

geschlossenen nach außen gekrümmten Fläche mit den Symmetrien des Würfels und zwölf, deren Schneiden mit den entsprechenden Würfelkanten verbindenden trapezartigen, ebenen Lamellen besteht (Fig. 19)***)). Dabei enthält jede geschlossene Blase $B^{(i)}$ ein ganz bestimmtes Luftquantum bei irgendeinem Volumen $V^{(i)}$ und irgendeinem Druck $p^{(i)}$, und ist demgemäß zur potentiellen Energie des gesamten Systems jedesmal noch der entsprechende Term $-p^{(i)}V^{(i)}$ hinzuzufügen. Dadurch folgt dann für eine Seitenfläche der Blase, auf welche auf der anderen Seite ein Druck $p^{(h)}$ herrscht, in Anbetracht der zwei Trennungsflächen der Lamellen anstatt

(17) allgemeiner und im Einklang mit (10):

$$p^{(i)} - p^{(h)} = 2T_{AB} \left(\frac{1}{R_1} + \frac{1}{R_2} \right),$$

*) H. A. Schwarz, Acta soc. scient. Fennicae 15 (1885), p. 315 (Ges. math. Abh. 1, S. 223).

**) Hilbert, Gött. Nachr. 1905, S. 159.

***) Plateau, l. c. p. 361.

wo die Krümmungsradien positiv bei nach außen konvexer Krümmung zu rechnen sind; zur Festlegung der $p^{(i)}$ dienen der Wert des äußeren Druckes sowie die Beträge der einzelnen eingeschlossenen Luftquanta. Die Randbedingungen beim Zusammentreffen dreier Flächen sind dieselben wie im früheren Falle nicht geschlossener Lamellen. So lassen sich z. B. mit Hilfe zweier fester Ringe wieder alle Formen des Unduloids und Nodoids, geschlossen durch angesetzte Kugelkalotten, erzielen. Eine einzelne freie Seifenblase hat notwendig Kugelgestalt und ist der Überdruck innen umgekehrt proportional ihrem Radius und der Proportionalitätsfaktor das Vierfache der Oberflächenspannung.

11. Stabilität einer Trennungsfläche.

Für das stabile Gleichgewicht einer Trennungsfläche F_{AB} , die bereits der früher erörterten Bedingung $\frac{dE}{dw} = 0 (w = 0)$ in jeder von einem Parameter w abhängenden und durch sie hindurch führenden Schar von virtuellen Verrückungen entspricht, ist weiter der definit-positive Charakter der zweiten Ableitung der potentiellen Energie nach dem Variationsparameter w , d. i. die Ungleichung:

$$(18) \quad \frac{d^2 E}{dw^2} > 0 \quad \text{für } w = 0$$

erforderlich. Nehmen wir an, der Rand von F_{AB} sei festzuhalten, so daß das Kurvenintegral in (3) fortfällt, so folgt durch Differentiation nach w aus (3), (4), (5) im Hinblick auf (6):

$$\frac{d^2 E}{dw^2} = \int_{F_{AB}} N \left\{ T_{AB} \frac{\partial}{\partial w} \left(\frac{1}{R_1} + \frac{1}{R_2} \right) + g(\varrho_A - \varrho_B) \frac{\partial z}{\partial w} \right\} df.$$

Hiervon sei eine spezielle Anwendung gemacht. Die Trennungsfläche falle in die Niveauebene $z = 0$. Die Flüssigkeit B befinde sich oberhalb A in einem nach unten offenen Gefäße, sei aber schwerer als A , also $\varrho_B > \varrho_A$ (Fig. 20). Hier ist für die variierten Flächen

$$z \equiv Nw \pmod{w^2}, \quad \frac{\partial}{\partial w} \left(\frac{1}{R_1} + \frac{1}{R_2} \right) \equiv -\frac{\partial^2 N}{\partial x^2} - \frac{\partial^2 N}{\partial y^2} \pmod{w},$$

(wobei durch das Zeichen $\equiv$ und den Zusatz $\pmod{w^2}$ bzw. $\pmod{w}$ eine Gleichheit bis auf Glieder von der Ordnung w^2 bzw. w angedeutet werden soll), und kommt die Bedingung (18) auf

$$(19) \quad \int_{F_{AB}} N \left\{ -T_{AB} \left(\frac{\partial^2 N}{\partial x^2} + \frac{\partial^2 N}{\partial y^2} \right) - g(\varrho_B - \varrho_A) N \right\} df > 0$$



Fig. 20.

hinaus, während die Konstanz des Volumens V_A die Gleichung

$$(20) \quad \int_{F_{AB}} N df = 0$$

erfordert und ferner N am Rande von F_{AB} durchweg Null sein soll.

In einer mehr elementaren Ausführung sagt Maxwell*), der Integrand in (19) müsse durchweg ≥ 0 sein, was überhaupt niemals für den ganzen Umfang der hier zuzulassenden Funktionen N zu erzielen wäre.

Ist die Öffnung des Gefäßes ein Kreis vom Radius R um den Nullpunkt, so trägt man der Bedingung des Verschwindens von $N(x, y)$ am Rande in allgemeinsten Weise Rechnung durch den Ansatz:

$$N(r \cos \varphi, r \sin \varphi) = \sum_{m=0}^{\infty} \sum_{k=1}^{\infty} J_m\left(\frac{\lambda_{mk} r}{R}\right) (a_{mk} \cos m\varphi + b_{mk} \sin m\varphi),$$

worin $J_m(\lambda)$ die Besselsche Funktion erster Art von der Ordnung m und $\lambda_{m1}, \lambda_{m2}, \dots$ ihre der Größe nach geordneten positiven Nullstellen bedeuten**). Aus (19) entsteht dann

$$\frac{\pi}{2} R^2 \sum_{m=0}^{\infty} \sum_{k=1}^{\infty} \left(\frac{m^2 \lambda_{mk}^2}{R^2} T_{AB} - g(\varrho_B - \varrho_A) \right) (J_{m+1}(\lambda_{mk}))^2 (a_{mk}^2 + b_{mk}^2) > 0,$$

während aus (20) die Gleichung:

$$2\pi R^2 \sum_{k=1}^{\infty} \frac{J_1(\lambda_{0k})}{\lambda_{0k}} a_{0k} = 0$$

wird. Nach der Größenfolge der λ_{mk} kommt die Forderung hier in der Tat auf das von Maxwell angegebene Kriterium für Stabilität hinaus, daß

$$R < \lambda_{11} \sqrt{\frac{T_{AB}}{g(\varrho_B - \varrho_A)}}$$

sein soll. Benutzt B die Gefäßwand, so kann die obere Grenze hier $\frac{\lambda_{11}}{\sqrt{2}} \sqrt{h_m}$ geschrieben werden, wenn h_m die mittlere Steighöhe in einer Kapillarröhre vom Radius 1 ist (vgl. Nr. 7)***); die Konstante $\lambda_{11}/\sqrt{2}$ hat den Wert 2,709...

Ist die Öffnung des Gefäßes ein Rechteck mit den Seiten a, b und $a \geq b$, so wird für die Stabilität des Gleichgewichtes

$$t^2 \left(\frac{4}{a^2} + \frac{1}{b^2} \right) T_{AB} - g(\varrho_B - \varrho_A) > 0$$

*) J. C. Maxwell, Scientific papers 2, p. 585.

**) Vgl. Die part. Differentialgl. d. math. Physik, nach Riemanns Vorl. neu bearbeitet von H. Weber, 2, S. 262; 1, S. 164.

***) Beobachtungen von Duprez (Mém. de l'Acad. de Belgique 26 (1851), 28 (1854)) sind in Übereinstimmung mit diesem theoretischen Ergebnisse.

erfordert, woraus zugleich für $a = \infty$ die entsprechende Bedingung in bezug auf einen langen Spalt ersichtlich ist.

12. Kapillarschwingungen.

Im Gleichgewichtszustand befinde sich A ganz unterhalb, B ganz oberhalb der Niveauebene $z = 0$, und ihre *unbegrenzt* gedachte Trennungsfläche führe nunmehr unter Einfluß der Oberflächenspannung und der Schwere flache Schwingungen

$$z \equiv \varepsilon f(x, y, t) \pmod{\varepsilon^2}$$

aus, worin ε einen Parameter in gewisser Umgebung von 0 bedeutet. In A wie in B mögen Geschwindigkeitspotentiale $\equiv \varepsilon \varphi_A$ bzw. $\equiv \varepsilon \varphi_B \pmod{\varepsilon^2}$ gelten, welche der Laplaceschen Differentialgleichung genügen und deren *negativ* genommene Differentialquotienten nach den Koordinaten die bezüglichen Geschwindigkeitskomponenten darstellen. An der Trennungsfläche haben wir einerseits für A , andererseits für B erstens die kinematische Forderung einer zur Fläche tangentialen Relativgeschwindigkeit, zweitens für den dort geltenden Druck $\equiv p_0 + \varepsilon p_A$ bzw. $\equiv p_0 + \varepsilon p_B \pmod{\varepsilon^2}$ das Integral der lebendigen Kraft und bestimmt sich drittens die Druckdifferenz $\equiv \varepsilon(p_A - p_B) \pmod{\varepsilon^2}$ als Kapillardruck gemäß (10). Für $\lim \varepsilon = 0$, d. h. für unendlich flache Wellen werden diese Beziehungen:

$$-\frac{\partial f}{\partial t} = \frac{\partial \varphi_A}{\partial z} = \frac{\partial \varphi_B}{\partial z}, \quad \frac{p_A}{\varepsilon_A} = \frac{\partial \varphi_A}{\partial t} - gf, \quad \frac{p_B}{\varepsilon_B} = \frac{\partial \varphi_B}{\partial t} - gf, \quad (z = 0),$$

$$p_A - p_B = -T_{AB} \left(\frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2} \right).$$

Soll noch A für $\lim z = -\infty$, B für $\lim z = +\infty$ ruhen, so wird allen hier genannten Bedingungen in einer Weise, die zur additiven Konstruktion ihrer allgemeinen Auflösung hinreicht, durch den partikulären Ansatz:

$$f = \Re(e^{-i\sigma t} F(x, y)), \quad \varphi_A = \Re\left(-\frac{i\sigma}{k} e^{kz-i\sigma t} F\right), \quad \varphi_B = \Re\left(\frac{i\sigma}{k} e^{-kz-i\sigma t} F\right),$$

$$\frac{\partial^2 F}{\partial x^2} + \frac{\partial^2 F}{\partial y^2} + k^2 F = 0$$

genügt, worin σ, k reelle positive Konstanten sind, $\Re$ das Zeichen für den reellen Teil der dahinter aufgeführten Größe ist und wo dann noch aus den letzten Relationen die Beziehung:

$$(21) \quad \left(\frac{\sigma}{k}\right)^2 = \frac{\varepsilon_A - \varepsilon_B}{\varepsilon_A + \varepsilon_B} \frac{g}{k} + \frac{T_{AB}}{\varepsilon_A + \varepsilon_B} k$$

zwischen k und σ folgt.

Wellen, die von y nicht abhängen, folgen bei dem Ansatz $F = C e^{ik(x-x_0)}$ und damit ist $\lambda = \frac{2\pi}{k}$ die Länge horizontal-zylindrischer, in einer Richtung

fortschreitender oder auch stehender Wellen von der Schwingungszahl $\frac{\sigma}{2\pi}$. Diese Beziehung (21) haben Lord Kelvin*), ferner Kolaček**) gegeben; sie findet Anwendung auf die Fortpflanzung von Wellen einer unbegrenzten

Wasserfläche unter der gemeinsamen Wirkung von Schwere und Kapillarität ohne Wind, ferner auf solche erzwungene stehende Kapillarschwingungen, bei denen die Knotenlinien als parallele Geraden gelten können***).

Der Gleichung (21) zufolge hat, $\varrho_A > \varrho_B$ vorausgesetzt, die Fortpflanzungsgeschwindigkeit $c = \frac{\sigma}{k}$ ein Minimum c_m bei einer gewissen Wellenlänge λ_m , mit welchen Größen dann (21) sich

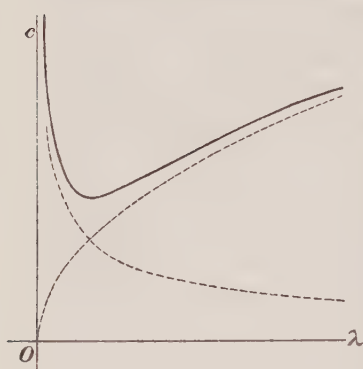


Fig. 21.

$$(21a) \quad \frac{c^2}{c_m^2} = \frac{1}{2} \left(\frac{\lambda}{\lambda_m} + \frac{\lambda_m}{\lambda} \right)$$

schreibt. (In Fig. 21 ist die hierdurch bestimmte Kurve in λ und c nebst den Kurven $c^2/c_m^2 = \frac{1}{2} \lambda/\lambda_m$ und $c^2/c_m^2 = \frac{1}{2} \lambda_m/\lambda$ dargestellt, um die Wirkungen von Schwere und Kapillarität zu vergleichen.) Mit einem jeden Werte $c > c_m$ vertragen sich alsdann zweierlei Wellenlängen, eine kürzere $\lambda_1 < \lambda_m$ und eine längere $\lambda_2 > \lambda_m$, wobei die Quotienten $\frac{\lambda_1}{\lambda_m}$ und $\frac{\lambda_2}{\lambda_m}$ reziprok sind. Die Wellen mit $\lambda < \lambda_m$, bei denen in (21) der Term mit T_{AB} gegenüber demjenigen mit g überwiegt, bezeichnet Lord Kelvin als „ripples“. Für Wasserwellen in Luft ist etwa $\lambda_m = 1,75$ cm, $c_m = 23,2 \frac{\text{cm}}{\text{sec}}$.

Lord Kelvin erörterte ferner den Einfluß des Windes auf die Geschwindigkeit von Wasserwellen. Hierbei wird die Annahme gemacht, daß die obere Flüssigkeit B für $\lim z = \infty$ mit einer gegebenen Geschwindigkeit u in Richtung der x -Achse fortschreitet. Bei der Wellenlänge λ sind alsdann zweierlei Fortpflanzungsgeschwindigkeiten

$$c_u = \frac{\varrho}{1 + \varrho} u \pm \sqrt{c^2 - \frac{\varrho}{(1 + \varrho)^2} u^2}, \quad \left(\varrho = \frac{\varrho_B}{\varrho_A} \right)$$

möglich, wo c die durch (21a) bestimmte Geschwindigkeit für $u = 0$ ist. Ein imaginärer Wert der Quadratwurzel hier würde bedeuten, daß die unverändert als Ausgangspunkt zu nehmende komplexe Partikulärlösung

*) W. Thomson, Phil. Mag. (4) 42 (1871), p. 368; Edinburgh Proc. Roy. Soc. 1870/71, p. 374.

**) Kolaček, Ann. Phys. Chem. 5 (1878), S. 425; 6 (1879), S. 616.

***) Vgl. die ausgedehnten Versuchsreihen von L. Grunmach, Wiss. Abh. d. kais. Normaleichungskommission, Berlin 1902, S. 101.

nunmehr in ihrem reellen Teile Wellen mit beständig zunehmender Amplitude darstellt. Diese Instabilität kommt für sämtliche Wellenlängen nicht in Frage, sowie $u < \frac{1+e}{e^{\frac{1}{2}}} c_m$ ist.

Denkt man sich wieder A in horizontal-zyklindrischer, von y nicht abhängender Bewegung derart, daß das von $z = 0$ wenig abweichende, im übrigen aber völlig willkürlich angesetzte Wellenprofil von A gleichförmig mit der Geschwindigkeit c in der x -Richtung fortschreitet und andererseits A für $z = -\infty$ ruht, so gewinnt man durch das Integral der lebendigen Kraft an der Oberfläche von A und andererseits den Kapillardruck eine Integralgleichung (Fouriersches Integral), um das Wellenprofil gerade einer willkürlich angenommenen Verteilung des äußeren Druckes p_B an der Oberfläche anzupassen. Insbesondere wirkt eine mit einer Geschwindigkeit $c > c_m$ in der x -Richtung schwimmende zur y -Achse parallele Gerade, welche an ihrem Orte den Gesamtbetrag des Druckes auf die Längeneinheit um P vermehrt, während sonst der Druck p_B konstant sei, genau wie eine sprungweise Zunahme des Richtungskoeffizienten dz/dx des Wellenprofils um den Betrag $2P/T_{AB}$ und ruft in einiger Entfernung vor sich her einfach-harmonische Wellen von der Länge $\lambda_1 (< \lambda_m)$, hinter sich von der Länge $\lambda_2 (> \lambda_m)$ hervor. — Eine gegen ihre Fortschreitungsrichtung einen Winkel $\frac{\pi}{2} - \theta$ bildende Drucklinie wirkt dann, als wenn sie nur senkrecht gegen sich die Geschwindigkeit $c \cos \theta$ hat, woraus durch eine Integration nach θ sich die Wirkung eines gleichförmig mit der Geschwindigkeit c schwimmenden, druckvermehrend wirkenden Punktes berechnet und insbesondere sich zeigt, daß ein solcher eine keilförmige Wellenfront (man denke an das Bild von Schiffswellen) mit dem durch $c \cos \theta = c_m$ bestimmten Öffnungswinkel $2\left(\frac{\pi}{2} - \theta\right)$ vor sich hertreibt*).

Die Berücksichtigung der *inneren Reibung* wird für flache, in einer Richtung fortschreitende Wellen auf einer *reinen* Wasseroberfläche derart zu geschehen haben, daß an der Oberfläche die Schubspannung gleich Null angenommen, die Zugspannung dem Kapillardruck entsprechend berechnet wird. Ist μ der Reibungskoeffizient und $\nu = \mu/\rho_A$, so findet man zu gegebener Wellenlänge λ anstatt der früheren Fortpflanzungsgeschwindigkeit c , wofern $\vartheta = 2\pi\nu/c\lambda$ klein ausfällt, (für Wasserwellen ist $\frac{2\pi\nu}{c_m} = 0,0048$ cm), eine modifizierte Wellengeschwindigkeit $= c(1 - \sqrt{2}\vartheta^{\frac{3}{2}})$, während zugleich die Amplituden einen Dämpfungsfaktor $e^{-\frac{8\pi^2\nu}{\lambda^2}t}$, also eine Relaxationszeit

*) Lord Rayleigh, Proc. Lond. Math. Soc. 15 (1883), p. 69 (Sc. papers 2, p. 258).

aufweisen*).
 $\frac{\lambda^2}{8\pi^2\nu}$ ($= 0,712 \lambda^2 \text{ sec}$ für Wasser)

Die beruhigende Einwirkung von Öl auf Wasserwellen wird dadurch erklärt**)*), daß zunächst infolge Überwiegens der Oberflächenspannung von Wasser gegen Luft über die Summe der zwei Oberflächenspannungen von Öl gegen Wasser und gegen Luft das Öl sich zu einer äußerst dünnen Haut auf dem Wasser auszieht und für die Oberflächenschicht mit der Beimengung von Öl dann elastische Eigenschaften zutage treten; ihre Spannung bleibt nicht länger konstant, sondern wächst, wenn die Dicke durch Streckung weiter zu reduzieren gesucht wird; dadurch wirkt sie gleichsam wie eine biegsame und schwer dehbare Membran und hindert durch ihren Zug auf das darunter befindliche Wasser die freie Entfaltung und Fortpflanzung der Wellen. Infolgedessen ist, wenn man den Einfluß der inneren Reibung ermitteln will, nicht mehr, wie im Falle einer reinen Wasseroberfläche, mit der Grenzbedingung an der Oberfläche zu rechnen, daß dort die Schubspannung Null ist, sondern eher mit der anderen, daß dort die horizontale Geschwindigkeitskomponente Null sei†). Für diesen anderen extremen Fall ergibt sich eine gegen die vorhin betrachteten Umstände im Verhältnis $4\sqrt{2}\vartheta : 1$ kleinere Relaxationszeit.

Die kleinen Schwingungen einer Trennungsfläche von der Gestalt eines *Kreiszyinders* behandelte Lord Rayleigh††), um von da aus die Stabilität der Flüssigkeitsstrahlen beurteilen zu können. Die Schwere wird nicht berücksichtigt. Es sei A innerhalb, B außerhalb des Zylinders befindlich, R der Radius des Zylinders, seine Achse die z -Achse, und

$$r \equiv R + \varepsilon f(z, \theta, t) \pmod{\varepsilon^2}, \quad (x + iy = r e^{i\theta}),$$

im $\lim \varepsilon = 0$ seine Schwingungsgleichung. Wird der äußere Druck in B konstant angenommen, was damit gleichwertig ist, $q_B = 0$ zu nehmen, so kann man für das Geschwindigkeitspotential in A den partikulären Ansatz

$$\varepsilon \Re(C e^{i(kz - \sigma t + m\theta)} J_m(ikr)) \pmod{\varepsilon^2}$$

machen, wobei J_m die Besselsche Funktion erster Art von der Ordnung m bedeutet, und man gelangt durch die kinematische Bedingung und andererseits die Druckgleichung an der Oberfläche zu der Relation

$$\sigma^2 = \frac{ikR J'_m(ikR)}{J_m(ikR)} (k^2 R^2 + m^2 - 1) \frac{T_{AB}}{\varrho_A R^3}.$$

*) Vgl. H. Lamb, *Hydrodynamics*, 3rd ed., Cambridge 1906, p. 563. [Deutsche Ausgabe, bes. v. J. Friedel, Leipzig 1907, S. 699.]

**) Reynolds, *Brit. Assoc. Rep.* 1880 (Sc. papers 1, p. 409).

***) Aitken, *Edinburgh Roy. Soc. Proc.* 12 (1883), p. 56.

†) H. Lamb, *Hydrodynamics*, 3rd ed., Cambridge 1906, p. 570. [Deutsche Ausg., S. 709.]

††) Lord Rayleigh, *Lond. Proc. Math. Soc.* 10 (1878), p. 4; *Proc. Roy. Soc.* 29 (1879), p. 71 (Sc. papers 1, p. 361, 377; *Theory of sound*, 2nd ed. chapt. XX); in

Für $m = 0$ wird $\sigma^2 < 0$, falls $kR < 1$ ist, was den instabilen Charakter von Störungen bedeutet, deren Wellenlänge $2\pi/k$ den Umfang des Zylinders überschreitet. Die Instabilität wird infolge des Faktors $e^{|\sigma|t}$ in den Amplituden am größten, wenn dabei $|\sigma|$ am größten ausfällt, was auf $\frac{2\pi}{k} = 4,51 \times 2R$ hinführt, so daß für Schwellungen und Kontraktionen von dieser Wellenlänge die Tendenz des Strahls A zum Zerfallen in Tropfen am stärksten ist.

Nach ähnlichen Prinzipien behandelt Lord Rayleigh*) den Fall $q_A = 0$, $q_B > 0$, wobei sich als die Wellenlänge größter Instabilität $\frac{2\pi}{k} = 6,48 \times 2R$ ergibt.

Das erste Ergebnis findet Anwendung auf das Zerfallen eines Wasserstrahls in Luft, das zweite auf das Zerreißen eines durch Wasser geschickten Luftstrahls. Die Schwingungen für $m = 2, 3, 4$ treten prädominierend hervor, wenn der Strahl aus einer Öffnung von elliptischer, dreieckiger, quadratischer Form austritt.

Die kleinen Schwingungen einer Trennungsfläche von der Gestalt einer Kugel erledigen sich ausgehend von dem gleichzeitigen Ansatz**)***)

$$\varphi_A = \Re \left(-\frac{C}{m} \frac{r^m}{R^m} Y_m(\theta, \psi) e^{-i\sigma t} \right), \quad \varphi_B = \Re \left(\frac{C}{m+1} \frac{R^{m+1}}{r^{m+1}} Y_m(\theta, \psi) e^{-i\sigma t} \right),$$

wobei der kinematischen Bedingung $\frac{d\varphi_A}{dr} = \frac{d\varphi_B}{dr}$ an der Oberfläche Rechnung getragen ist; darin bedeuten r, θ, ψ Polarkoordinaten vom Kugelmittelpunkt, $Y_m(\theta, \psi)$ die Kugelflächenfunktion m^{ter} Ordnung, R den Radius der Kugel. Es stellt sich alsdann

$$\sigma^2 = m(m+1)(m-1)(m+2) \frac{T_{AB}}{((m+1)q_A + mq_B)R^3}$$

heraus. Das Ergebnis findet Anwendung auf die Schwingungen eines Wassertropfens in Luft, einer Luftblase in Wasser; in abfallenden Tropfen treten durch ein Nachwirken des Abreißens der Tropfen noch die Schwingungen 3. Ordnung ($m = 3$) hervor†).

Phil. Mag. 34 (1892), p. 145 (Sc. papers 3, p. 585) wird noch der Einfluß der inneren Reibung der Flüssigkeit in Betracht gezogen.

*) Lord Rayleigh, Phil. Mag. (5) 34 (1892), p. 177 (Sc. papers 3, p. 594).

**) Lord Rayleigh, Proc. Roy. Soc. 29 (1879), p. 71 (Sc. papers 1, p. 377).

***) Webb, Mess. of math. 9 (1880), p. 177.

†) Lenard, Ann. Phys. Chem. 30 (1887), S. 209.

II. Kapillarität als räumlich verteilte Energie.

13. Die Hypothese der Kohäsionskräfte.

Die Kapillaritätserscheinungen ergeben sich als notwendige Folgerungen aus einer Hypothese, wonach zwischen zwei materiellen Teilchen gleicher oder verschiedener Substanzen neben der Gravitation noch eine andere, nur von der Distanz abhängende Anziehungskraft in der Verbindungslinie wirksam ist, die man *Kohäsionskraft* nennt und deren Gesetz irgendwelcher Art sein mag, nur daß sie mit wachsender Entfernung derart rasch abnimmt, daß sie bereits auf eine äußerst kleine, mikroskopisch nicht wahrnehmbare Distanz ganz außer Betracht fällt.

Zunächst wurde das Ansteigen von Flüssigkeit in einer kapillaren Röhre allein mit einer von der Röhre auf die Flüssigkeit ausgeübten Anziehung erklärt, die nach der Unabhängigkeit der Erscheinung von der Dicke der Röhre nur von den der Wand nächstgelegenen Partikeln ausgehen konnte*). Clairaut**) erkannte es als notwendig, eine Anziehung der Flüssigkeitsteilchen untereinander mit in Rücksicht zu ziehen. Laplace***) konnte sodann eine vollständige Theorie der Kapillarität einzig mit der vorhin skizzierten Hypothese über die Kohäsionskräfte aufbauen.

Laplace berechnete für eine Flüssigkeitsmasse, deren Teile gemäß jener Hypothese kohärieren, in der Hauptsache das Potential der Kohäsionskräfte für eine Stelle der Oberfläche und fand es als eine lineare Funktion der mittleren Krümmung daselbst. Er betrachtete zunächst das Potential einer Kugel auf eine Stelle der Oberfläche, ging von da zum Potential eines durch zwei unendlich nahe Meridianschnitte der Kugel gebildeten Keiles über und approximierte endlich eine beliebige Flüssigkeitsoberfläche in der Nähe eines Punktes durch die dort aus den Krümmungskreisen der Normalschnitte erzeugte Fläche, d. i. ungefähr durch das Oskulationsparaboloid. Die Differentialgleichung einer freien Oberfläche erhielt er nunmehr aus dem Satze der Hydrostatik, wonach diese bei konstantem äußeren Drucke eine Fläche konstanten Potentials aller wirkenden Kräfte ist†).

*) Hawkesbee, London Trans. R. Soc. 26, 27 (1709—1713).

**) Clairaut, *Traité sur la figure de la terre*, Paris 1743, chap. X.

***) Laplace, *Théorie de l'action capillaire*.

†) Genauer gesagt, verfuhr Laplace so: er dachte sich in der Flüssigkeit einen unendlich schmalen Kanal gelegt, der am Anfang und Ende senkrecht gegen die Oberfläche einmündet, berechnete den durch die Kohäsionskräfte zustande kommenden Druck auf einen Querschnitt des Kanals und brachte endlich das „Prinzip des Gleichgewichts in Kanälen“ zur Anwendung.

In einer zweiten Darstellung berechnete Laplace*) für eine Stelle der Flüssigkeitsoberfläche die tangentielle Komponente der gesamten dort ausgeübten Kohäsionskraft, wozu die Flächengleichung an der Stelle bis einschließlich der Größen 3^{ter} Ordnung zu entwickeln ist, und erhielt die Gleichung der freien Oberfläche aus der Bedingung, daß an ihr die Resultante aus Kohäsion und Schwere stets normal zur Fläche steht. — Für die Konstanz des Randwinkels der Flüssigkeit an einem festen Körper hatte jedoch Laplace keinen Beweis, sondern zeigte nur, daß, wenn der Körper die Form von vertikalen Zylindern irgendwelcher Querschnitte hat, der Mittelwert des Kosinus jenes Winkels längs der ganzen Randkurve stets auf die nämliche Konstante führen muß.

Diese Lücke in der Laplaceschen Theorie ergänzte Gauß**). Ausgehend von dem Prinzip der virtuellen Verrückungen für einen Gleichgewichtszustand formte Gauß dieses Prinzip zu der Forderung eines Minimums der potentiellen Energie um und betrachtete nunmehr die gesamte potentielle Energie der ins Spiel kommenden Kohäsionskräfte. Diese Energie erscheint in Form eines Doppel-Raumintegrals. Für jede Raumintegration läßt sich eine lineare ausführen, wobei ein Term proportional dem Volumen und ein zweiter proportional dem Flächeninhalt der Oberfläche besonders heraustreten, und von dem übrig bleibenden Doppel-Oberflächenintegral wies Gauß nach, daß bei der Laplaceschen Annahme über das Abnehmen der Kohäsionskraft die Vernachlässigung desselben geboten ist, insofern als verglichen mit der Distanz, auf die allein die Kohäsionskräfte in Betracht kommen, die Krümmungsradien der Oberfläche als unendlich groß, die Fläche als nahezu eben gelten darf. Aus dem extremalen Charakter der potentiellen Energie entnahm sodann Gauß nach den Methoden der Variationsrechnung (ungefähr wie oben in Nr. 3 und 4 dargelegt ist) die Differentialgleichung für die freie Oberfläche, dann aber auch den Beweis für den Laplaceschen Satz vom konstanten Randwinkel.

14. Potentielle Energie der Kohäsion in einem Medium.

Die von Gauß vorgenommene Transformation der Energie von Kohäsionskräften***) läßt sich gegenwärtig als eine zweimalige Anwendung des Greenschen Satzes darstellen.

*) Laplace, Suppl. à la théorie de l'act. capill.

**) Gauß, Principia generalia, Göttingen 1830. — Selbstanzeige der Abh.: Gött. gel. Anz. 1829, (Werke 5, S. 287).

***) Vereinfachte Darstellungen dieser Transformation gaben Bertrand, Journ. de math. (1) 13 (1848), p. 185; Weinstein, Ann. Phys. Chem. 27 (1886), S. 544. — L. Boltzmann, Ann. Phys. Chem. 141 (1870), S. 582 setzte Summationen über Moleküle an Stelle der Integrationen.

Wir betrachten zunächst die Kohäsionsenergie innerhalb eines einzelnen homogenen Mediums A . Seine Dichte heiße ϱ ; zwischen je zwei Volumenelementen dv, dv' von A an den Stellen x, y, z und x', y', z' im Abstände r wirke als Kohäsion eine Anziehungskraft $= \varrho^2 dv dv' \varphi(r)$, wo $\varphi(r)$ in später noch genau festzusetzender Weise mit zunehmendem r rapide nach Null sinken soll. Führt man

$$\int_r^\infty \varphi(r) dr = \psi(r), \quad \int_r^\infty r^2 \psi(r) dr = \chi(r)$$

ein, so läßt sich die gesamte potentielle Energie dieser Kohäsionskräfte für A :

$$(22) \quad -\frac{1}{2} \varrho^2 \iint \psi(r) dv dv' \\ = -\frac{1}{2} \varrho^2 \iint \left(\frac{\partial \chi}{\partial x'} \frac{\partial 1}{\partial x'} + \frac{\partial \chi}{\partial y'} \frac{\partial 1}{\partial y'} + \frac{\partial \chi}{\partial z'} \frac{\partial 1}{\partial z'} \right) dv dv'$$

schreiben, wobei einerseits dv , andererseits dv' unabhängig voneinander das ganze Volumen von A zu durchlaufen hat und der Faktor $\frac{1}{2}$ vorzusetzen ist, weil auf diese Weise jedes Paar von Elementen dv, dv' zweimal in Betracht kommt. Fassen wir zunächst die Integration nach dv' bei festgehaltenem dv ins Auge, so können wir den zweiten Ausdruck in (22) nach dem Greenschen Satze umformen, wobei wir die Laplacesche Differentialgleichung für $1/r$ verwenden, aber infolge der Unstetigkeit von $1/r$ an der Stelle dv noch aus dem Integrationsraume für dv' eine kleine Kugel um dv auszuschneiden haben, deren Radius wir schließlich nach Null konvergieren lassen. Bezeichnen wir mit df'' (später auch mit df) ein Oberflächenelement von A , mit n' (und n) die äußeren Normalen dort, so transformiert sich nunmehr (22) in:

$$-2\pi \varrho^2 \chi(0) \int dv - \frac{1}{2} \varrho^2 \int dv \int \chi(r) \frac{\partial 1}{\partial n'} df''.$$

Im ersten Term hier ist $\int dv = V_A$ das gesamte Volumen von A . Im zweiten Term kehren wir die Integrationsfolgen um, führen

$$\int_r^\infty \chi(r) dr = \vartheta(r)$$

ein und beachten, daß r nur von den Differenzen $x - x', y - y', z - z'$ abhängt, ferner

$$\frac{\partial}{\partial n'} \left(\left(\frac{\partial r}{\partial x} \right)^2 + \left(\frac{\partial r}{\partial y} \right)^2 + \left(\frac{\partial r}{\partial z} \right)^2 \right) = 0$$

ist. Wir können alsdann diesen zweiten Term

$$\frac{1}{2} \varrho^2 \int df' \int \left(\frac{\partial \vartheta(r)}{\partial x} \frac{\partial r}{\partial n'} \frac{\partial \frac{1}{r}}{\partial x} + \frac{\partial \vartheta(r)}{\partial y} \frac{\partial r}{\partial n'} \frac{\partial \frac{1}{r}}{\partial y} + \frac{\partial \vartheta(r)}{\partial z} \frac{\partial r}{\partial n'} \frac{\partial \frac{1}{r}}{\partial z} \right) dv$$

schreiben und erhalten die Möglichkeit, ein zweites Mal bei der Integration nach dv die Formel für Produktintegration, den Greenschen Satz, in Anwendung zu bringen. Hier liegt die Unstetigkeit von $1/r$ jedesmal an einer Oberflächenstelle df' und ist deshalb aus dem Integrationsraume für dv nur der in den Bereich von A fallende Teil einer kleinen Kugel um diese Stelle auszuschneiden, d. i. hernach bei unendlich abnehmendem Radius der Kugel wesentlich eine Halbkugel, außer an den Stellen, wo eine Schneide der Oberfläche von A vorhanden ist. Hiernach transformiert sich der letzte Ausdruck, indem noch $\int \frac{\partial r}{\partial n'} df$ über die Halbkugel ihre Projektion auf die Tangentialebene in df' darstellt, in

$$\frac{1}{2} \pi \varrho^2 \vartheta(0) \int df' - \frac{1}{2} \varrho^2 \iint \frac{1}{r^2} \frac{\partial r}{\partial n} \frac{\partial r}{\partial n'} \vartheta(r) df df',$$

wo $\int df' = F$ den Flächeninhalt der Oberfläche von A bildet. Schreiben wir noch

$$(23) \quad K = 2\pi \varrho^2 \chi(0), \quad H = \pi \varrho^2 \vartheta(0),$$

so resultiert endlich der folgende Ausdruck für die Energie der Kohäsionskräfte innerhalb A :

$$(24) \quad E = -KV + \frac{1}{2} HF - \frac{1}{2} \varrho^2 \iint \frac{1}{r^2} \frac{\partial r}{\partial n} \frac{\partial r}{\partial n'} \vartheta(r) df df'.$$

Damit die Integrale für $\vartheta(r)$, $\chi(r)$, $\psi(r)$ einen Sinn haben, nehmen wir an, daß mit wachsendem r jedenfalls $r\chi(r)$, sodann $r^4\psi(r)$, $r^5\varphi(r)$ noch hinreichend stark nach Null konvergieren. Des weiteren nehmen wir an, daß für mikroskopisch meßbare r schon $\varphi(r)$, $\psi(r)$, $\chi(r)$, $\vartheta(r)$ außer Betracht fallen und erst bei weit stärkerer Annäherung des r an Null $\chi(r)$ und $\vartheta(r)$ endliche Größe erlangen und dann bestimmten oberen Grenzen $\chi(0)$ und $\vartheta(0)$ zustreben; es ist hierfür notwendig und hinreichend, daß $r^3\psi(r)$ für $\lim r=0$ nach Null konvergiert. Bezeichnet man als *Wirkungsradius* für die Kohäsionskräfte eine solche Größe r_0 , wofür eben noch $\vartheta(r_0)$ gegen $\vartheta(0)$ zu vernachlässigen ist, so ersieht man aus

$$\vartheta(0) - \vartheta(r_0) = \int_0^{r_0} \chi(r) dr < \chi(0)r_0,$$

daß $\frac{\vartheta(0)}{\chi(0)}$ eine äußerst kleine Länge, allenfalls von der Ordnung von r_0 sein wird.

Um das Doppel-Flächenintegral in (24) abzuschätzen, führen wir dort durch $\frac{df'}{r^2} \frac{\partial r}{\partial n'} = d\omega'$ den körperlichen Winkel ein, unter dem das Element

df' von der Stelle von df erscheint. Jenes Integral schreibt sich dann

$$(25) \quad -\frac{1}{2}\varrho^2 \int df \int \frac{\partial r}{\partial n} \vartheta(r) d\sigma'.$$

Nun ist nur für kleine $r (< r_0)$ der Faktor $\vartheta(r)$ von merklicher Größe, und für solche $r < r_0$ andererseits ist $\partial r / \partial n$ angenähert $= r/R$, unter R den Krümmungsradius desjenigen Normalschnittes für die Stelle df , der durch df' führt, verstanden. Sind also die Krümmungsradien der Oberfläche von A überall als gegen r_0 äußerst groß zu betrachten, so erscheint die Vernachlässigung des Integrals hier geboten und reduziert sich damit der Ausdruck (24) auf

$$(26) \quad E = -KV + \frac{1}{2}HF.$$

Ein Ausnahmefall wird statthaben, wenn die Oberfläche A eine Partie $\mathfrak{S}$ von endlicher Ausdehnung aufweist, zu der eine andere Partie $\mathfrak{S}'$ von ihr fortwährend in äußerst kleinen Abständen verläuft, wie z. B. wenn die Flüssigkeit sich in einer äußerst dünnen Schicht an einem festen Körper entlang zieht. Auch längs $\mathfrak{S}$ und $\mathfrak{S}'$ mögen die Krümmungsradien nicht unter eine gegen r_0 äußerst große Grenze sinken. Faßt man in dem Integral (25) ein Element df innerhalb $\mathfrak{S}$ mit der ganzen Partie $\mathfrak{S}'$ zusammen auf und fällt ein Lot von df auf $\mathfrak{S}'$, dessen Länge s sei, so kann in denselben Fehlergrenzen, wie sie vorhin galten, $\mathfrak{S}'$ als eine unbegrenzt ausgedehnte Ebene senkrecht auf dieses Lot betrachtet und, indem man $\frac{\partial r}{\partial n'} = \frac{s}{r} = \cos \gamma$ einführt, aus konzentrischen Ringen um das Lot, die von df in körperlichen Winkeln $2\pi \sin \gamma d\gamma$ erscheinen, aufgebaut werden; dazu kann wegen der annähernden Parallelität von df zu $\mathfrak{S}'$ der Faktor $\partial r / \partial n$ in (25) durch $\partial r / \partial n'$ ersetzt werden, wodurch für den betreffenden Anteil aus (25) sich ergibt

$$-\frac{1}{2}\varrho^2 df \int_0^{\frac{\pi}{2}} 2\pi \sin \gamma \cos \gamma \vartheta(r) d\gamma = -\pi \varrho^2 df \int_s^{\infty} \frac{s^2}{r^3} \vartheta(r) dr.$$

Wird nun

$$\theta(r) = 2r^2 \int_r^{\infty} \frac{\vartheta(r)}{r^3} dr = \vartheta(r) + r^2 \int_r^{\infty} \frac{\chi(r)}{r^2} dr$$

eingeführt, wobei $\theta(0) = \vartheta(0)$ ersichtlich ist, und beachtet man, daß im Doppelintegral (25) sowohl eine Kombination $\mathfrak{S}, \mathfrak{S}'$ wie $\mathfrak{S}', \mathfrak{S}$ auftritt, so kommt schließlich wegen der Flüssigkeitsschicht zwischen $\mathfrak{S}$ und $\mathfrak{S}'$ zum Ausdruck (26) der Energie noch der Zusatzterm

$$(27) \quad -\pi \varrho^2 \int_{\mathfrak{S}} \theta(s) df,$$

über die ganze eine Seite $\mathfrak{S}$ der Schicht erstreckt, hinzu.

Für das Anziehungsgesetz mit folgendem Ausdrucke des Potentials*)

$$(28) \quad -\psi(r) = -k \frac{e^{-cr}}{r},$$

wobei k und c positive Konstanten sind, würde man erhalten:

$$\varphi(r) = k \frac{e^{-cr}}{r^2} (1 + cr),$$

$$\chi(r) = \frac{k}{c^2} e^{-cr} (1 + cr), \quad \vartheta(r) = \frac{2k}{c^3} e^{-cr} \left(1 + \frac{1}{2} cr\right), \quad \theta(r) = \frac{2k}{c^3} e^{-cr},$$

$$\chi(0) = \frac{k}{c^2}, \quad \vartheta(0) = \theta(0) = \frac{2k}{c^3}, \quad c = \frac{2\chi(0)}{\vartheta(0)} = \frac{K}{H}.$$

Wir berechnen noch das *Virial der Kohäsionskräfte*. (Unter dem Virial einer Kraft wird bekanntlich die halbe Arbeit der Kraft bei Verschiebung ihres Angriffspunktes nach dem Koordinatenanfange verstanden.) Für die zwei Kräfte, welche zwei Volumenelemente dv, dv' gegenseitig aufeinander ausüben, ist die Summe der Viriale $q^2 \varphi(r) r dv dv'$. Das gesamte Virial der Flüssigkeit auf sich selbst wird daher

$$\frac{1}{4} q^2 \iint r \varphi(r) dv dv'$$

und würde aus dem Ausdrucke (22) für die Energie hervorgehen, wenn man $\psi(r)$ dort durch $-\frac{1}{2} r \varphi(r)$ ersetzt. Dabei würde dann an Stelle von $\chi(r)$ die Funktion

$$-\frac{1}{2} \int_r^\infty r^3 \varphi(r) dr = -\frac{1}{2} r^3 \psi(r) - \frac{3}{2} \int_r^\infty r^2 \psi(r) dr = -\frac{1}{2} r^3 \psi(r) - \frac{3}{2} \chi(r),$$

weiter an Stelle von $\vartheta(r)$ die Funktion

$$-\frac{1}{2} \int_r^\infty r^3 \psi(r) dr - \frac{3}{2} \int_r^\infty \chi(r) dr = -\frac{1}{2} r \chi(r) - 2 \vartheta(r)$$

zu treten haben und also die Rolle der Konstanten $\chi(0), \vartheta(0)$ von $-\frac{3}{2} \chi(0), -2 \vartheta(0)$ übernommen werden. Der Formel (26) entsprechend resultiert dann als Ausdruck jenes Gesamtvirials:

$$(29) \quad \frac{3}{2} K V - H F.$$

15. Potentielle Energie der Adhäsion zweier Medien.

Grenzt A an ein zweites Medium B , so mögen zwischen den Teilchen von A und denen von B Anziehungskräfte wirksam sein, die hinsichtlich ihrer Abnahme mit der Distanz analogen Charakter tragen wie die Ko-

*) Van der Waals, Zeitschr. f. physik. Chem. 13 (1894), S. 657.

häsionskräfte innerhalb A , und die man hier im Falle verschiedener Substanzen wohl auch als *Adhäsionskräfte* bezeichnet. Die den Funktionen $\varphi(r)$, $\psi(r)$, $\chi(r)$, $\vartheta(r)$, $\theta(r)$ oben entsprechenden Funktionen für das neue Anziehungsgesetz mögen in derselben Weise unter Anfügung der unteren Indizes AB bezeichnet werden, während die früheren Funktionen den Index A erhalten mögen. Die gesamte Energie der Adhäsion von B auf A berechnet sich der Formel (22) analog mit

$$(30) \quad -\varrho_A \varrho_B \int_B dv' \int_A \psi_{AB}(r) dv,$$

wo dv die Volumenelemente von A , dv' diejenigen von B zu durchlaufen hat. Ein Faktor $\frac{1}{2}$ ist jetzt nicht hinzuzusetzen, weil die Räume von A und B völlig getrennt sind. Der Ausdruck hier gestattet wieder die zwei entsprechenden Umformungen mittels des Greenschen Satzes. Aber in der ersten Umformung, wobei etwa unter Festhaltung der einzelnen dv operiert werde, tritt kein Raumintegral besonders heraus, weil jetzt $1/r$ im Integrationsraume B keine Unstetigkeitsstelle hat, in der zweiten hernach kommt bei festgehaltenem Oberflächenelement df' von B eine Diskontinuität von $1/r$ im Integrationsraume A nur für solche Elemente df' in Betracht, die gerade der Trennungsfläche von A und B angehören, und hat alsdann für die um df' herum aus dem Integrationsraume A auszuscheidende kleine Halbkugel das Integral $\int \frac{\partial r}{\partial n} df$ wegen der anders liegenden Normale n' entgegengesetzten Wert wie oben. Infolge dieser Umstände erlangt endlich nach der wie oben vorzunehmenden Vernachlässigung die potentielle Energie der Adhäsion von B auf A den Ausdruck

$$(31) \quad -\pi \varrho_A \varrho_B \vartheta_{AB}(0) F_{AB} = -H_{AB} F_{AB},$$

wo F_{AB} den Flächeninhalt der Trennungsfläche von A und B bezeichnet.

Nehmen wir jetzt den typischen Fall von drei zusammentreffenden Medien A , B , C an und bezeichnen ihre Volumina V , ihre Konstanten K und H mit den entsprechenden einzelnen Indizes, die Flächeninhalte ihrer Trennungsflächen und deren Konstanten H mit entsprechenden zwei Indizes, während ihre an weitere Medien angrenzenden Flächen hier nicht als veränderlich in Frage kommen sollen, so haben wir als veränderlichen Teil der potentiellen Energie der in ihnen wirkenden Anziehungskräfte:

$$(32) \quad \left(\frac{1}{2} H_A + \frac{1}{2} H_B - H_{AB}\right) F_{AB} + \left(\frac{1}{2} H_A + \frac{1}{2} H_C - H_{AC}\right) F_{AC} \\ + \left(\frac{1}{2} H_B + \frac{1}{2} H_C - H_{BC}\right) F_{BC} - K_A V_A - K_B V_B - K_C V_C.$$

So lange die Volumina sich nicht ändern, kommen wir damit auf den Ansatz in Nr. 2 zurück, wobei die Oberflächenspannung zweier Medien

A und B durch

$$(33) \quad T_{AB} = \frac{1}{2} H_A + \frac{1}{2} H_B - H_{AB}$$

gegeben erscheint. Im Falle $\varrho_B = 0$ gesetzt werden kann, folgt einfach

$$T_{AB} = \frac{1}{2} H_A.$$

Stellt C einen festen Körper vor und darf $\varrho_B = 0$ gesetzt werden, so folgt für den Randwinkel ω_A von A am Körper gemäß Gl. (8):

$$(34) \quad \cos \omega_A = \frac{2H_{AC} - H_A}{H_A};$$

der Winkel ω_A ist spitz oder stumpf und es findet demgemäß, falls C ein vertikaler Zylinder ist, ein Ansteigen oder eine Depression von A an C hinsichtlich der Niveauebene statt, je nachdem $2H_{AC} > H_A$ oder $< H_A$ ist (d. h. wenn man so sagen will, der Meniskus vom Körper eine mehr oder weniger als doppelt so starke Anziehung wie von der Flüssigkeit erfährt*).

Die Relation (34) ist unmöglich, wenn $H_{AC} > H_A$ ist. Nehmen wir jedoch an, daß alsdann sich die Flüssigkeit A noch in einer äußerst dünnen Schicht an einer Partie $\mathfrak{S}$ der Wand von C entlang zieht, und verstehen wir jetzt unter F_{AB} , F_{AC} nur die Flächeninhalte der betreffenden Trennungsflächen abgesehen von dieser Schicht, so würde im Hinblick auf (27) und auf die Relation $\vartheta(0) = \theta(0)$ zum Ausdrucke (32) in der gesamten Energie noch ein Term hinzutreten:

$$-\int_{\mathfrak{S}} \left(H_{AC} \left(1 - \frac{\theta_{AC}(s)}{\theta_{AC}(0)} \right) - H_A \left(1 - \frac{\theta_A(s)}{\theta_A(0)} \right) \right) df,$$

erstreckt über die Fläche $\mathfrak{S}$, wobei s die Dicke der Schicht am Elemente df von $\mathfrak{S}$ bezeichnet. Bei geeignetem Kraftgesetz, z. B. dem in (28) angeführten, würde damit die Möglichkeit einer Verringerung der potentiellen Energie vermöge der Schicht vorliegen; also müßte dann eine solche Schicht (Benetzung der Wand) zustande kommen und dadurch am Rande der wahrnehmbaren Trennungsfläche der Randwinkel Null entstehen**).

16. Eingehen der Kohäsion in die Beziehung zwischen Dichte und Druck.

Das Auftreten des Terms $-KV$ in der Energie einer Flüssigkeit A ist nach hydrodynamischen Prinzipien gleichbedeutend mit der Annahme, daß im Inneren von A neben dem sogenannten hydrostatischen Druck ein weiterer konstanter Druck K herrscht. Schreibt man $K = a\varrho_A^2$, so hängt a

*) Clairaut, *Traité sur la figure de la terre*, Paris 1743, chap. X.

**) Gauß, *Principia generalia*, art. 32.

nur von dem Kraftgesetz $\varphi(r)$, nicht von der Dichte ϱ_A ab. Stellt man sich unter B den gesättigten Dampf der Flüssigkeit A vor, so hat daher eine Menge M der Substanz —, homogene kontinuierliche Massenverteilung bis zu den Grenzen und Unabhängigkeit des Kohäsionsgesetzes von der Temperatur angenommen (wegen allgemeinerer Vorstellungen siehe den Enzyklopädie-Artikel V 10 von Kamerlingh Onnes) und vorausgesetzt auch, daß nicht etwa F/V gegen $1/r_0$ in Betracht kommt —, in flüssiger Phase die Energie $-K \frac{M}{\varrho_A} = -a\varrho_A M$, in dampfförmiger die Energie $-a\varrho_B \frac{M}{\varrho_B} = -a\varrho_B M$ und wird deshalb $a(\varrho_A - \varrho_B)$ als *die innere latente spezifische Verdampfungswärme* (siehe ebenfalls Enzyklopädie V 10) angesprochen*). (Die additive Konstante in der Energie war derart fixiert, daß der Wert Null für die Energie als obere Grenze bei Auflösung des Mediums in lauter unendlich weit voneinander entfernte Volumenelemente entsteht.)

In denjenigen Erscheinungen, welche bei konstantem Volumen vorgehen, kommt die Größe K gar nicht zur Geltung, während doch der erste Term $-KV$ in der Energie den anderen $\frac{1}{2}HF$ außerordentlich überwiegt. Aufschluß über die Größe von K kann deshalb nur von Vorgängen erwartet werden, die mit Änderungen der Dichte verknüpft sind, und van der Waals**) hatte deshalb die Idee, zunächst theoretisch das Eingehen dieser Größe in die Beziehung zwischen Druck und Dichte bei konstanter Temperatur zu untersuchen. Der Ableitung dieser Beziehung legte van der Waals den Virialsatz von Clausius zugrunde***). Der Satz kommt bei folgenden Anschauungen zustande: Die Materie ist nicht kontinuierlich verteilt, sondern besteht aus Molekülen; diese unterliegen neben den Laplaceschen Kohäsionskräften weiteren repulsiven Kräften (Zusammenstoßen) und sind dadurch in nicht sichtbaren Bewegungen begriffen, und zwar dergestalt, daß man bei Vergleichung ihrer Bewegungszustände in irgend zwei Momenten t und $t + \tau$ sich angenähert vorstellen kann, die Teilchen hätten nur untereinander Ort und Bewegung gewechselt. Jedenfalls soll, wenn man den Mittelwert von der kinetischen Energie der progressiven Bewegung der Moleküle über den Zeitraum t bis $t + \tau$ bildet, gegen diesen Mittelwert der Differenzenquotient

$$\frac{1}{\tau} \left[\frac{1}{4} \frac{d \Sigma m(x^2 + y^2 + z^2)}{dt} \right]_t^{t+\tau}$$

schon bei verhältnismäßig kleinem τ zu vernachlässigen sein; darin ist die

*) Dupré, Théorie mécanique de la chaleur, 1869, p. 152.

**) Van der Waals, Die Kontinuität des gasf. u. flüss. Zustandes, Leipzig 1881.

***) Vgl. auch Maxwell, Sc. papers 2, p. 407, 418; H. A. Lorentz, Boltzmann-Festschrift (1904), S. 721 (abgedr. in Abh. üb. theor. Phys. 1 (1906), S. 192).

Summe über alle Moleküle zu erstrecken und bedeuten m die Masse, x, y, z die Koordinaten des Schwerpunktes eines Moleküls. Eine partielle Integration transformiert nun diesen Mittelwert der Energie bei der angegebenen Vernachlässigung sofort in das mittlere Virial der in den Molekülen angreifenden Kräfte, über den Zeitraum t bis $t + \tau$. Nun ist nach den Prinzipien der Gastheorie jenes Mittel der Energie der progressiven Bewegung $\frac{3}{2} \frac{R}{M} \varrho VT$, wo R die universelle Gaskonstante, M das Molekulargewicht, ϱV die Gesamtmasse, T die absolute Temperatur der Flüssigkeit vorstellt. Das Virial der Kohäsionskräfte berechnet sich unter der Voraussetzung, daß der Wirkungsradius noch sehr groß gegen die Größe der Moleküle ist, wie bei kontinuierlich homogen verteilter Masse und ist daher nach (29) gleich $\frac{3}{2} a \varrho^2 V$ zu setzen. Das mittlere Virial des auf die Oberfläche wirkenden konstanten Druckes p findet sich durch Zerlegung des Volumens in die Elementarpyramiden mit dem Nullpunkte als Spitze und den Oberflächenelementen als Grundflächen unmittelbar $= \frac{3}{2} p V$. Das mittlere Virial der repulsiven Kräfte berechnet sich nach den Methoden der Gastheorie und läßt sich schreiben als ein Bruchteil $\frac{b \varrho}{1 - b \varrho}$ der mittleren Energie der progressiven Bewegung, worin b annäherungsweise als konstant gilt und mit dem von den Räumen der Moleküle in einer Masseneinheit besetzten Raume in Verbindung gebracht wird (über die Abhängigkeit der Größe b von Volumen und Temperatur siehe weiteres im Artikel Kamerlingh Onnes, Enzyklopädie V 10). Damit resultiert endlich die van der Waalssche Zustandsgleichung in der Form

$$(35) \quad p + a \varrho^2 = \frac{R T}{M} \frac{\varrho}{1 - b \varrho}.$$

Aus beobachteten Daten berechnet sich auf Grund dieser Beziehung bei 0° und 1 Atm. Druck für Wasser $K = 10500$ Atm., für Äther $K = 1430$ Atm., während der Quotient H/K , der als Maß für den Wirkungsradius der Kohäsionskräfte dient, für Wasser $15 \cdot 10^{-9}$ cm, für Äther $29 \cdot 10^{-9}$ cm beträgt.

17. Theorien zur Vermeidung von Diskontinuitäten der Dichte.

Der Laplaceschen Kapillaritätstheorie liegt die Annahme durchweg homogener Flüssigkeiten zugrunde. Poisson*) führte aus, daß an den Grenzflächen einer Flüssigkeit eine rapide Änderung der Dichte statthaben

*) Poisson, Nouvelle théorie de l'action capillaire. — Kritische Bemerkungen zur Theorie von Poisson lieferten Minding, Doves Repert. d. Phys. Bd. 5; J. Stahl, Ann. Phys. Chem. 139 (1870), S. 239; B. Weinstein, Ann. Phys. Chem. 27 (1886), S. 544.

müsse, und trug diesem Umstande Rechnung, zugleich in der Absicht, die Schwierigkeiten zu beheben, welche in der Annahme von Druckdiskontinuitäten an Trennungsflächen liegen. In der Poissonschen Theorie modifizieren sich nicht die Gleichungen für die Kapillaritätsphänomene, sondern nur die Bedeutung der zwei Konstanten K und H für das Gesetz der Kohäsionskräfte.

Maxwell*), Lord Rayleigh**), van der Waals***) verfolgten die Annahme einer stetigen Variation der Dichte an den Trennungsflächen in ihren weiteren Konsequenzen. Als einfachster Fall wird das Gleichgewicht einer Flüssigkeit A in Berührung mit ihrem gesättigten Dampfe B behandelt. Von der Schwere soll abgesehen werden. Der ganze Raum von Flüssigkeit und Dampf werde von Flächen, auf denen jedesmal die Dichtigkeit konstant ist, durchzogen; transversal zu diesen variiert der Wert der Dichtigkeit rapide innerhalb einer äußerst schmalen Schicht und kommt nach der einen wie nach der anderen Seite sehr bald bestimmten Grenzwerten ϱ_A bzw. ϱ_B nahe.

Für die vollständige Durchführung des durch hydrodynamische (oder thermodynamische) Prinzipien gelieferten Ansatzes hat sich ein spezielles Kraftgesetz der Kohäsion als hervorragend geeignet erwiesen†), das nun hier sogleich zugrunde gelegt werde, nämlich das in (28) angeführte, wobei die Potentialfunktion für zwei Masseneinheiten in einer Entfernung r durch

$$-k \frac{e^{-cr}}{r}$$

dargestellt wird und k wie c positive Konstanten sind. Das von der gesamten Masse der Substanz (A und B) herrührende Potential auf eine Masseneinheit an einer Stelle x, y, z ist dann

$$\Psi(x, y, z) = -k \int \varrho' \frac{e^{-cr}}{r} dv',$$

wo das Integral sich über alle Volumenelemente dv' der Substanz erstreckt und r die Entfernung des Aufpunktes x, y, z vom Elemente dv' bezeichnet. Diese Funktion $\Psi(x, y, z)$ genügt nun im Raume der Substanz überall der Differentialgleichung

$$(36) \quad \Delta \Psi = \frac{\partial^2 \Psi}{\partial x^2} + \frac{\partial^2 \Psi}{\partial y^2} + \frac{\partial^2 \Psi}{\partial z^2} = c^2 \Psi + 4\pi k \varrho,$$

und auf Grund dieser Beziehung kann das vorliegende Kraftfeld anstatt

*) Maxwell, Capillary action (Sc. papers 2, p. 541).

**) Lord Rayleigh, Phil. Mag. 33 (1892), p. 209 (Sc. Papers 3, p. 513).

***) van der Waals, Zeitschr. f. phys. Chemie 13 (1894), S. 657; H. Hulshof, Ann. Phys. Chem. (4) 4 (1901), S. 165.

†) van der Waals, l. c. S. 706.

als von Fernkräften herstammend auch als ein ursprünglich gegebener Spannungszustand in der Substanz folgendermaßen beschrieben werden*). Es werde

$$\left(\frac{\partial \Psi}{\partial x}\right)^2 + \left(\frac{\partial \Psi}{\partial y}\right)^2 + \left(\frac{\partial \Psi}{\partial z}\right)^2 = \Phi^2,$$

$$\frac{1}{8\pi k}(c^2\Psi^2 - \Phi^2) = \Sigma_1, \quad \frac{1}{8\pi k}(c^2\Psi^2 + \Phi^2) = \Sigma_2$$

gesetzt; die Substanz erscheint von den gerichteten „Kraftlinien“, die senkrecht zu den Flächen $\Psi = \text{konst.}$ von größeren zu kleineren Werten von Ψ führen, durchzogen; an jeder Stelle herrscht in Richtung der dort durchlaufenden Kraftlinie und in der entgegengesetzten Richtung eine Zugspannung Σ_1 , in allen Richtungen senkrecht dazu eine Zugspannung Σ_2 , dergestalt, daß für jeden abgeschlossenen Teil I der Substanz sich die Komponenten und Drehungsmomente der von dem übrigen Teil II auf I ausgeübten Kohäsionen genau wie aus der Verteilung dieser Spannungen auf der Oberfläche von I berechnen.

In einer sehr geringen Entfernung von der Übergangsschicht ist bereits nahezu Ψ konstant, Φ Null, und $\Sigma_1 = \Sigma_2$ erklären die frühere Kohäsion K , in der Übergangsschicht resultieren aus der Differenz $\Sigma_2 - \Sigma_1$ die Erscheinungen der Oberflächenspannung.

Nach hydrodynamischen Prinzipien ist für das Gleichgewicht des Systems Flüssigkeit und Dampf bei gleicher Temperatur erforderlich, daß das vollständige Differential

$$(37) \quad d\Psi = -\frac{d\Pi}{\varrho}$$

und darin Π eine nur von Dichte und Temperatur der Stelle abhängige Funktion ist, die als *thermischer Druck***) angesprochen wird. Schreiben wir $\frac{2\pi k}{c^2} = a$, so ist nach (36) in einiger Entfernung von der Übergangsschicht, wo die Flüssigkeit homogen erscheint, $\Psi = -2a\varrho_A$, und wo der Dampf homogen erscheint, $\Psi = -2a\varrho_B$. Wir setzen allgemein $\Pi = p + a\varrho^2$ und nennen p den *hydrostatischen Druck*; nach beiden Seiten von der Schicht fort wird dann p sich einer und derselben Konstante p_0 nähern, dem äußeren Drucke, Sättigungsdrucke des Dampfes. Die Gleichung (36) schreibt sich noch im Hinblick auf (37):

$$(38) \quad \Delta\Psi = -c^2 \int_{p_0, \varrho_A}^{p, \varrho} \frac{dp}{\varrho}.$$

*) G. Bakker, Zeitschr. f. phys. Chemie 48 (1904), S. 17.

**) Diese Bezeichnung gebraucht van der Waals; für denselben Begriff sagt H. A. Lorentz (Z. f. physik. Chem. 7): kinetischer Druck.

Nunmehr wird angenommen, daß die Abhängigkeit des p von ϱ und der Temperatur auch in der Übergangsschicht durch eben dieselbe van der Waalssche Formel (35) wie in den homogenen Phasen dargestellt wird, was freilich mehr an die Macht dieser Formel glauben heißt, als es in ihrer Ableitung eine Stütze fände. Die dadurch gegebene Kurve für p als Funktion des wachsenden Arguments $1/\varrho$ (siehe Artikel Kamerlingh Onnes, Enzyklopädie V 10) verläuft in dem Intervalle $1/\varrho_A$ bis $1/\varrho_B$ zwischen den zwei Punkten $p_0, 1/\varrho_A$ und $p_0, 1/\varrho_B$ ab-, auf- und wieder absteigend, zuerst unterhalb der Geraden $p = p_0$ bis zu einem gewissen Treffpunkte mit ihr, hernach oberhalb derselben, und auf ihrem ersten unterhalb $p = p_0$ liegenden Stücke muß es offenbar einen bestimmten Punkt $p_1, 1/\varrho_1$ geben, für den das bis dahin auf der Kurve genommene Integral

$$-\int_{p_0, \varrho_A}^{p_1, \varrho_1} \frac{dp}{\varrho} = 0$$

ist (Fig. 22). Für diese Stelle der Kurve ist dann nach (38): $\Delta\Psi = 0$.

Die der Wellenlinie ($\varrho_A > \varrho > \varrho_B$) entsprechenden Kombinationen $p, 1/\varrho$ sind für homogene Phasen instabil, in der Übergangsschicht findet van der Waals sie nun stabil. Wird (37) auf einem Wege aus dem homogenen Inneren der Flüssigkeit nach dem homogenen Inneren des Dampfes integriert, so entsteht

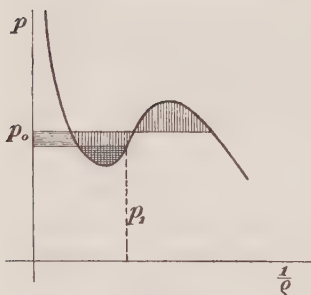


Fig. 22.

$$\int_{p_0, \varrho_A}^{p_0, \varrho_B} \frac{dp}{\varrho} = 0$$

über jene Wellenlinie, genau die Formel, welche Clausius und Maxwell zur Bestimmung des Druckes p_0 des gesättigten Dampfes vermöge der Isotherme durch Anwendung des zweiten Hauptsatzes der Wärmetheorie auf labile Zustände aufgestellt haben.

Es mögen nun die Flächen konstanter Dichte speziell als eine Schar paralleler Ebenen $z = \text{konst.}$ angenommen werden. Die Differentialgleichung (37) schreibt sich dann

$$\frac{d^2\Psi}{dz^2} = c^2\Psi + 4\pi k\varrho = c^2\Psi - 4\pi k \frac{d\Pi}{d\Psi}$$

und liefert integriert

$$\left(\frac{d\Psi}{dz}\right)^2 = c^2\Psi^2 - 8\pi k(p + a\varrho^2 - p_0).$$

Für $\varrho = \varrho_1, p = p_1$ geht $\frac{d^2\Psi}{dz^2}$ abnehmend durch Null hindurch, erlangt

also $\frac{d\Psi}{dz} = \Phi(z)$ seinen größten Wert $\Phi_1 = \sqrt{8\pi k(p_0 - p_1)}$; hierhin werde $z = 0$ gelegt. An Stelle der früheren Oberflächenspannung tritt jetzt

$$\frac{1}{2}\bar{H} = \int (\Sigma_2 - \Sigma_1) dz = \frac{1}{4\pi k} \int (\Phi(z))^2 dz,$$

auf der z -Achse aus der homogenen Flüssigkeit nach dem homogenen Dampf genommen.

Die Kurve für $\Phi(z)$ als Funktion von z verläuft beiderseits von $z = 0$ sehr bald asymptotisch an die z -Achse. Wird sie durch das oberhalb der z -Achse liegende Stück einer sie im Scheitel berührenden Parabel $\Phi_1 - \Phi = \Phi_1 \left(\frac{z}{z_0}\right)^2$ und die links und rechts daran ansetzenden Stücke $z \leq -z_0$ und $z_0 \leq z$ der z -Achse bei gleichem Flächeninhalt über der z -Achse angenähert ersetzt, so folgt

$$\frac{1}{2}\bar{H} = \frac{8}{5} \frac{\sqrt{2\pi k(p_0 - p_1)}}{c^3},$$

während zugleich $2z_0 = \frac{15}{16} \frac{\frac{1}{2}\bar{H}}{p_0 - p_1}$ ungefähr als diejenige Dicke der Übergangsschicht aufzufassen wäre, innerhalb deren die inhomogenen Verhältnisse hervortreten.

In einer jüngsten Arbeit sucht Bakker*) durch seine Theorie die Beobachtungen von Reinold und Rücker**) zu erklären, welche fanden, daß Seifenlamellen an den dünnsten Stellen, die durch ein diskontinuierliches Auftreten von schwarzen Flecken gekennzeichnet sind, eine Dicke von rund 10^{-6} cm haben und unmittelbar daneben plötzlich auf eine Dicke von rund $5 \cdot 10^{-6}$ cm ansteigen.

Bakker berechnet daselbst die Konstante k

für Wasser $k = 7,53 \cdot 10^{23}$ bei $T = 325^\circ$,

für Äther $k = 1,54 \cdot 10^{23}$ bei $T = 125^\circ$.

18. Entropie und Massendichten einer Trennungsfläche.

Wenn zwei aneinander grenzende Flüssigkeiten A und B sich im Gleichgewicht, auch in thermischer und chemischer Hinsicht, befinden und sie auch schon in sehr geringer Entfernung von der Trennungsfläche homogen erscheinen, wird doch eine jede in der unmittelbaren Nachbarschaft der Grenze durch den Einfluß der anderen verändert sein. Gibbs***) hat einen Ansatz entwickelt, um diesen Einflüssen Rechnung zu tragen, ohne

*) G. Bakker, Zeitschr. f. phys. Chemie 51 (1905), S. 344.

**) Proc. Roy. Soc. 26 (1878), p. 334; Phil. Trans. 172 (1882), p. 447; Phil. Trans. 174 (1884), p. 645; Ann. Phys. Chem. 44 (1891), S. 778.

***) Gibbs, Equilibrium of heterogeneous substances, p. 380.

irgendeine Hypothese bezüglich molekularer Anziehungskräfte zu machen. Die inhomogene Übergangsschicht zwischen A und B ist erfahrungsgemäß von äußerst geringer Dicke. Man wähle irgendeinen Punkt in dieser Schicht und lege eine Fläche durch ihn und alle anderen Punkte in der Schicht, welche hinsichtlich der unmittelbar angrenzenden Materie entsprechend liegen; diese Fläche heiße *Teilungsfläche*. Die Wahl der Fläche ist noch einigermaßen willkürlich; man wird annehmen können, daß man sie beliebig aus einer Schar sehr nahe gelegener Parallellflächen, welche die ganze Schicht ausfüllen, herausgreifen kann. Die in A , B und in der Übergangsschicht in Betracht kommenden Stoffverbindungen mögen sich aus den Stoffen a , b , ... als *unabhängigen* Bestandteilen aufbauen lassen. Für das ganze aus A , B und der Übergangsschicht bestehende Gebilde sei U die gesamte innere Energie, S die gesamte Entropie und seien M_a , M_b , ... die gesamten Massen von a , b , ... Wir bezeichnen mit V' , V'' die Volumina von A und B , gerechnet bis zur Teilungsfläche, mit F den Flächeninhalt der Teilungsfläche. Weiter seien u' , s' , ϱ_a' , ϱ_b' , ... die räumlichen Dichten der Energie, Entropie bzw. diejenigen der Bestandteile a , b , ... im Raume von A dort, wo A homogen erscheint, und u'' , s'' , ϱ_a'' , ϱ_b'' , ... die entsprechenden Dichten für B , wo B homogen erscheint. Die Quotienten aus den Differenzen

$$U - V'u' - V''u'', S - V's' - V''s'', M_a - V'\varrho_a' - V''\varrho_a'', \dots$$

durch den Flächeninhalt F endlich schreiben wir u , s , ω_a , ω_b , ...; diese Quotienten heißen die *Flächendichten der Energie, Entropie und der Massenbestandteile* für die Teilungsfläche zwischen A und B .

Es wird angenommen, daß u' eine Funktion der Argumente s' , ϱ_a' , ϱ_b' , ..., desgleichen u'' eine Funktion der Argumente s'' , ϱ_a'' , ϱ_b'' , ... ist, und nunmehr die weitere Annahme eingeführt, daß auch u nur eine Funktion der Argumente s , ω_a , ω_b , ... ist (vgl. hierzu die am Anfange von Nr. 5 berührte allgemeine Auffassung der räumlichen Energiedichte), und es wird daraufhin die Charakterisierung des Gleichgewichtszustandes durch das thermodynamische Prinzip geleistet, daß U ein Minimum bei konstanten Werten von S , M_a , M_b , ... ist. Dieser Ansatz, soweit er die homogenen Massen betrifft, ist bereits im Artikel Bryan, Enzyklopädie V 3 Nr. 26, zur Sprache gekommen. Hier soll es sich nur darum handeln, die besonderen Konsequenzen, welche aus der neuen Annahme einer Übergangsschicht fließen, zu verfolgen.

Entwickelt man das vollständige Differential

$$(39) \quad du = Tds + \mu_a d\omega_a + \mu_b d\omega_b + \dots,$$

so ist T die Temperatur und μ_a , μ_b , ... heißen die *Potentiale* für die Bestandteile der Übergangsschicht. Zum Gleichgewicht ist erforderlich,

daß die Temperatur dieselbe wie in den angrenzenden homogenen Massen ist, weiter daß für jeden Bestandteil der Schicht, der auch in einer angrenzenden homogenen Masse vorkommt, das Potential das gleiche wie dort ist. Kommt ein Bestandteil nur in der Schicht vor, so ist für ihn dagegen die Flächendichte in der Schicht von vornherein angewiesen.

Aus (39) folgt

$$d(Fu) = Td(Fs) + \sigma dF + \mu_a d(F\omega_a) + \mu_b d(F\omega_b) + \dots,$$

indem

$$(40) \quad \sigma = u - Ts - \mu_a \omega_a - \mu_b \omega_b - \dots$$

gesetzt wird; σ ist nunmehr als die *Oberflächenspannung* in der Teilungsfläche anzusprechen, und es entsteht

$$(41) \quad d\sigma = -s dT - \omega_a d\mu_a - \omega_b d\mu_b - \dots$$

Eine Beziehung, welche u als Funktion von $s, \omega_a, \omega_b, \dots$ oder σ als Funktion von $\omega_a, \omega_b, \dots$ gibt, wird als *Fundamentalgleichung* für die Teilungsfläche bezeichnet.

Analog hat man in der homogenen Masse A :

$$(42) \quad d(V'u) = Td(V's) - p'dV' + \mu_a d(V'\varrho_a) + \mu_b d(V'\varrho_b) + \dots, \\ -p' = u' - Ts' - \mu_a \varrho_a' - \mu_b \varrho_b' - \dots,$$

(und zwar mit den nämlichen Werten von T und von $\mu_a, \mu_b, \dots$, soweit die betreffenden Bestandteile in A wirklich vorkommen), und die entsprechenden Beziehungen in der homogenen Masse B ; dabei bedeuten dann p' und p'' den hydrostatischen Druck in A bzw. in B . Für die Gestalt der Teilungsfläche folgt, wie (10) gewonnen wurde:

$$(43) \quad p' - p'' = \sigma \left(\frac{1}{R_1} + \frac{1}{R_2} \right).$$

Von der Schwere ist einstweilen abgesehen. Die Dichten $u, s, \omega_a, \omega_b, \dots$ sind im allgemeinen noch abhängig von der Wahl der Teilungsfläche in der Schicht; im Falle einer ebenen Teilungsfläche ist jedoch der Wert σ von dieser Wahl unabhängig, wie sich leicht auf Grund von (43) ergibt.

Sind die unabhängigen Bestandteile $a, b, \dots$ derart fixiert, daß nicht $\varrho_a' < 0$ sein kann und wird die homogene Phase A nur in bezug auf ϱ_a' bei konstant gehaltenen $T, p', \varrho_b', \dots$ abgeändert, so spricht das Verhalten der idealen Gase und verdünnten Lösungen dafür, daß $\varrho_a' \frac{d\mu_a}{d\varrho_a'}$ mit nach Null abnehmendem ϱ_a' einem bestimmten (und aus Stabilitätsrücksichten notwendig) positiven Grenzwerte zustrebt, daß mithin für sehr kleine Werte ϱ_a' das Potential μ_a wesentlich durch einen Ausdruck $K_a \log \varrho_a'$ dargestellt

werden kann, worin K_a eine positive Funktion von $T, p', \varrho_a', \dots$ bedeutet*). Die Gesamtmenge von a ist

$$M_a = V' \varrho_a' + V'' \varrho_a'' + F \omega_a;$$

wenn nun weder M_a noch ϱ_a', ϱ_a'' negativ sein können, so kann bei kleinen Werten der räumlichen Dichten ϱ_a', ϱ_a'' die Flächendichte ω_a beliebig große positive, aber nur kleine negative Werte haben. Die Relation (41) zeigt nunmehr auf Grund der eben gemachten Ausführungen, daß eine geringe Beimengung eines Stoffes die zwischen zwei Medien bestehende Oberflächenspannung sehr stark vermindern, aber nicht sie beträchtlich vermehren kann. Ganz frisch gebildete Oberflächen lassen in der Regel eine Änderung der Oberflächenspannung nicht erkennen, insofern die Herstellung der statischen Oberflächendichte Zeit beansprucht**).

Bringt man auf eine reine Wasseroberfläche kleine Kampherstückchen, so löst sich der Kampher an den Berührungsstellen mit dem Wasser und wird dort die Oberflächenspannung herabgesetzt. Die Kampherpartikelchen geraten in lebhafte Bewegung dadurch, daß diese Änderung der Oberflächenspannung an verschiedenen Stellen in verschiedenem Maße geschieht***). Lord Rayleigh†) fand, daß die Bewegung des Kamphers erlischt, wenn die Wasseroberfläche durch eine Ölschicht, gleichgültig welcher Art, soweit verunreinigt wird, bis die Spannung eben auf den Betrag sinkt, den sie für Wasser, das mit Kampher gesättigt ist, erlangt, nämlich bis auf 53 erg/cm², d. i. 72% des Wertes 74 für reines Wasser. Diesen „Kampherpunkt“ bringt bei Olivenöl eine Schicht von ungefähr 2×10^{-7} cm Dicke hervor. Lord Rayleigh††) hat den Einfluß noch geringerer Verunreinigungen des Wassers auf die Oberflächenspannung gemessen und fand, daß der Abfall der Spannung erst bei einer Ölschicht von etwa 1×10^{-7} cm mehr diskontinuierlich beginnt und nach Überschreitung des Kampherpunktes wieder träger wird.

Der Ansatz (39) gibt Gibbs insbesondere Gelegenheit zu mannigfachen Ausführungen über das Verhalten von Flüssigkeitshäuten.

Dieselben Prinzipien bringt Gibbs auch auf die Trennungsflächen einer Flüssigkeit gegen einen festen amorphen oder kristallinen Körper in Anwendung. Bei Kristallen würde die Spannung σ an den begrenzenden

*) Gibbs, l. c. p. 194.

**) A. Dupré, *Théorie mécanique de la chaleur*, Paris 1869, p. 377; Lord Rayleigh, *Proc. Roy. Soc.* 47 (1890), p. 281 (Sc. papers 3, p. 341).

***) Van der Mensbrugghe, *Mém. couronnés de l'Acad. de Belg.* 34 (1869).

†) Lord Rayleigh, *Phil. Mag.* 30 (1890) (Sc. papers 3, p. 383).

††) Lord Rayleigh, *Proc. Roy. Soc.* 47 (1890), p. 364 (Sc. papers 3, p. 347); A. Pockels, *Nature* 43 (1891), p. 437; 48 (1893), p. 153.

Flächen noch eine Funktion der Stellung der Flächen im Kristall sein, und die Gestalt sehr kleiner, im Gleichgewicht mit der umgebenden Lösung stehender Kristalle würde wesentlich durch die Bedingung bestimmt sein, daß die Summe $\sum \sigma F$ der Oberflächenenergien aller einzelnen Flächen ein Minimum in bezug auf das vorliegende Volumen ist.

Bei Rücksichtnahme auf die Schwere modifiziert sich infolge des Ansatzes (41) die frühere Bedingung (43) für die Gestalt der Trennungsfläche zu den Gleichungen:

$$p' - p'' = \sigma \left(\frac{1}{R_1} + \frac{1}{R_2} \right) + g\omega \cos(nz), \quad \frac{d\sigma}{dz} = g\omega,$$

wobei n die nach B hin gerichtete Normale, $\frac{1}{R_1} + \frac{1}{R_2}$ die mittlere Krümmung nach B hin und $\omega = \omega_a + \omega_b + \dots$ die totale Massendichtigkeit an der variablen Stelle der Teilungsfläche bedeutet.

Die thermischen Beziehungen, zu denen der Ansatz (41) hinführt, waren bereits zuvor von W. Thomson*) durch Betrachtung geeigneter Kreisprozesse gewonnen worden. Um ein Beispiel zu geben, befinde sich eine ebene flüssige Lamelle A in gesättigtem Dampfe B . Die Teilungsfläche sei auf jeder Seite der Lamelle derart gelegt, daß die Oberflächendichte des einzigen vorhandenen Bestandteiles a Null wird, dann folgt aus (41): $d\sigma = -s dT$. Wird die Lamelle auseinander gezogen und leitet man den Vorgang isothermisch und ohne Kondensation, also bei festbleibendem Gesamtvolumen, so ist für die Einheit entstandener Fläche die zuzuführende Wärmemenge

$$Q = Ts = -T \frac{d\sigma}{dT} = -\frac{d\sigma}{d \log T}.$$

Leitet man dagegen den Vorgang adiabatisch und bei konstantem äußeren Druck p , so gehört zu p als Druck des gesättigten Dampfes eine festbleibende Verdampfungstemperatur T , und um diese aufrecht zu erhalten, muß sich für die Einheit entstandener Fläche eine Menge δ_a Dampf kondensieren, deren Betrag sich durch die Bedingung der Wärmezufuhr Null:

$$s - \delta_a(s'' - s') = 0$$

ergibt. Dem entspricht eine Volumzunahme δ_a/ϱ_a' der Flüssigkeit und eine Volumabnahme δ_a/ϱ_a'' des Dampfes und ist danach zur Erhaltung des Druckes von außen her für die Einheit entstandener Oberfläche die Arbeit

$$W = ps \left[\frac{1}{s'' - s'} \left(\frac{1}{\varrho_a''} - \frac{1}{\varrho_a'} \right) \right]$$

*) W. Thomson, Proc. Roy. Soc. 9 (1858), p. 255 oder Phil. Mag. (4) 17 (1859), p. 61; van der Mensbrugghe, Bull. de l'Acad. de Bruxelles 51 (1876), p. 769, 52 (1876), p. 21; P. Duhem, Ann. de l'Éc. Norm. sup. (3) 2 (1885), p. 207.

zu leisten. Der hier in eckige Klammern gesetzte Ausdruck folgt aus der Abhängigkeit zwischen Temperatur und Druck des gesättigten Dampfes als der Differentialquotient dT/dp (vgl. Artikel Bryan, Enzyklopädie V 3, Gl. (138)) und es entsteht demnach

$$W = -p \frac{d\sigma}{dT} \frac{dT}{dp} = -\frac{d\sigma}{d \log p}.$$

Das Produkt aus σ in die $\frac{2}{3}$ -te Potenz des Molekularvolumens M_v der Flüssigkeit (d. i. des Volumens einer Menge M Grammen, wenn M das Molekulargewicht bezeichnet), nennt Ostwald *molekulare Oberflächenenergie* der Flüssigkeit. Der Differentialquotient $-d(M_v^{\frac{2}{3}}\sigma)/dT$, (man könnte ihn *molekulare Oberflächenentropie* nennen), liegt für eine große Reihe von Flüssigkeiten nahe am Werte $2,1 \text{ erg/cm}^2 \text{ grad}^*$). Dieses merkwürdige, von Eötvös**) und von Ramsay und Shields***) experimentell festgestellte Gesetz (siehe darüber den Artikel Kamerlingh Onnes, Enzyklopädie V 10) bietet eine Methode dar, das Molekulargewicht von Flüssigkeiten auf Grund von Kapillarkonstanten zu ermitteln.

Zu denjenigen Kapillaritätserscheinungen, deren Theorie sich auf den Gibbsschen Ansatz (41) basieren läßt, gehört endlich die Abhängigkeit der Oberflächenspannung einer flüssigen Metallelektrode gegen einen Elektrolyten von der elektromotorischen Kraft ihrer Polarisation†).

Inhaltsübersicht.

	Seite
1. Kapillarität und Kohäsion	299
I. Kapillarität als Flächenenergie.	
2. Oberflächenenergie und deren Variation	299
3. Differentialgleichung für eine freie Oberfläche	303
4. Randwinkel	305
5. Kapillardruck. Oberflächenspannung	309
6. Formen freier Oberflächen. Tropfen	311
7. Steighöhen.	315
8. Kapillarauftrieb. Adhäsion	316

*) Für Wasser, das hier zu den Ausnahmen gehört, liegt die molekulare Oberflächenentropie zwischen 0,9 und 1,2; danach ist für eine Wasserlamelle die oben besprochene latente Ausdehnungswärme nahezu gleich der halben zu ihrer Bildung gegen die Oberflächenspannung aufzuwendenden Arbeit ($Q = \frac{1}{2}\sigma$).

**) Eötvös, Ann. Phys. Chem. 27 (1886), S. 452.

***) Ramsay und Shields, Zeitschr. f. physik. Chem. 12 (1893), S. 433.

†) Für die Theorien in diesem experimentell reich durchforschten Kapitel der Elektrokapillarität vgl. F. Krüger, Gött. Nachr. 1904, S. 33; Jahrbuch d. Radioaktivität und Elektronik 2 (1904).

	Seite
9. Ausschaltung der Schwere	319
10. Flüssigkeitshäute	322
11. Stabilität einer Trennungsfläche	325
12. Kapillarschwingungen	327

II. Kapillarität als räumlich verteilte Energie.

13. Die Hypothese der Kohäsionskräfte	332
14. Potentielle Energie der Kohäsion in einem Medium	333
15. Potentielle Energie der Adhäsion zweier Medien	337
16. Eingehen der Kohäsion in die Beziehung zwischen Dichte und Druck	339
17. Theorien zur Vermeidung von Diskontinuitäten der Dichte	341
18. Entropie und Massendichten einer Trennungsfläche	345



Die Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern.

(Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen.

Mathematisch-physikalische Klasse. 1908. S. 53—111.)

(Vorgelegt in der Sitzung vom 21. Dezember 1907.)

Einleitung.

Über die Grundgleichungen der Elektrodynamik für bewegte Körper herrschen zur Zeit noch Meinungsverschiedenheiten. Die Ansätze von Hertz*) (1890) mußten verlassen werden, weil sich herausgestellt hat, daß sie mit verschiedenen experimentellen Ergebnissen in Widerspruch geraten.

1895 publizierte H. A. Lorentz**) seine Theorie der optischen und elektrischen Erscheinungen in bewegten Körpern, die, auf atomistischer Vorstellung von der Elektrizität fußend, durch ihre großen Erfolge die kühnen Hypothesen, von denen sie getragen und durchsetzt wird, zu rechtfertigen scheint. Die Lorentzsche Theorie***) geht aus von gewissen ursprünglichen Gleichungen, die an jedem Punkte des „Äthers“ gelten sollen, und gelangt daraus durch Mittelwertbildungen über „physikalisch unendlich kleine“ Bereiche, die schon zahlreiche „Elektronen“ enthalten, zu den Gleichungen für die Vorgänge in ponderablen Körpern.

Insbesondere gibt sich die Lorentzsche Theorie Rechenschaft von der Nichtexistenz einer Relativbewegung der Erde gegen den Lichtäther; sie bringt diese Tatsache in Zusammenhang mit einer Kovarianz jener ursprünglichen Gleichungen bei gewissen gleichzeitigen Transformationen der Raum- und Zeitparameter, die von H. Poincaré†) den Namen *Lorentz-Transformationen* erhalten haben. Für jene ursprünglichen Gleichungen ist die Kovarianz bei den Lorentz-Transformationen eine rein mathematische Tatsache, die ich das *Theorem der Relativität* nennen will; dieses Theorem beruht wesentlich auf der Gestalt der Differentialgleichung für die Fortpflanzung von Wellen mit Lichtgeschwindigkeit.

*) Über die Grundgleichungen der Elektrodynamik für bewegte Körper. Wiedemanns Annalen, Bd. 41, S. 369, 1890 (auch in: Gesammelte Werke, Bd. I, S. 256, Leipzig 1892).

**) Versuch einer Theorie der elektrischen und optischen Erscheinungen in bewegten Körpern, Leiden 1895.

***) Vgl. Enzyklopädie der mathematischen Wissenschaften, V 2, Art. 14. Weiterbildung der Maxwell'schen Theorie. Elektronentheorie.

†) Rendiconti del Circolo Matematico di Palermo, T. XXI (1906), p. 129.

Nun kann man, ohne noch zu bestimmten Hypothesen über den Zusammenhang von Elektrizität und Materie sich zu bekennen, erwarten, jenes mathematisch evidente Theorem werde seine Konsequenzen so weit erstrecken, daß dadurch auch die noch nicht erkannten Gesetze in bezug auf ponderable Körper irgendwie eine Kovarianz bei den Lorentz-Transformationen übernehmen werden. Man äußert damit mehr eine Zuversicht, als bereits eine fertige Einsicht, und diese Zuversicht will ich das *Postulat der Relativität* nennen. Die Sachlage ist erst ungefähr eine solche, als wenn man die Erhaltung der Energie postuliert in Fällen, wo die auftretenden Formen der Energie noch nicht erkannt sind.

Gelangt man hernach dazu, die erwartete Kovarianz als einen bestimmten Zusammenhang zwischen lauter beobachtbaren Größen bei bewegten Körpern zu behaupten, so mag alsdann dieser bestimmte Zusammenhang das *Prinzip der Relativität* heißen.

Diese Unterscheidungen scheinen mir nützlich, um den gegenwärtigen Stand der Elektrodynamik bewegter Körper charakterisieren zu können.

H. A. Lorentz hat das Relativitätstheorem gefunden und das Relativitätspostulat geschaffen, als eine Hypothese, daß Elektronen und Materie infolge von Bewegung Kontraktionen nach einem gewissen Gesetze erfahren.

A. Einstein*) hat es bisher am schärfsten zum Ausdruck gebracht, daß dieses *Postulat* nicht eine künstliche Hypothese ist, sondern vielmehr eine durch die Erscheinungen sich aufzwingende neuartige Auffassung des Zeitbegriffs.

Das *Prinzip* der Relativität jedoch in dem von mir gekennzeichneten Sinne ist für die Elektrodynamik bewegter Körper bisher noch gar nicht formuliert worden. *In der gegenwärtigen Abhandlung erhalte ich, indem ich dieses Prinzip formuliere, die Grundgleichungen für bewegte Körper in einer durch dieses Prinzip völlig eindeutig bestimmten Fassung. Dabei wird sich zeigen, daß keine der bisher für diese Gleichungen angenommenen Formen sich diesem Prinzip genau fügt.*

Man sollte vor allem erwarten, daß die von Lorentz angenommenen Grundgleichungen für bewegte Körper dem Relativitätspostulate entsprächen. Es zeigt sich indes, daß dieses nicht der Fall ist für die allgemeinen Gleichungen, die Lorentz für beliebige, auch magnetisierte Körper hat, daß es aber allerdings *approximativ* (unter Vernachlässigung der Quadrate der Geschwindigkeiten der Materie gegen das Quadrat der Lichtgeschwindigkeit) der Fall ist für diejenigen Gleichungen, die Lorentz

*) Annalen der Physik, Bd. 17, S. 891, 1905.

hernach für nichtmagnetisierte Körper erschließt; es kommt aber diese spätere Anpassung an das Relativitätspostulat wieder nur dadurch zustande, daß die Bedingung des Nichtmagnetisiertseins ihrerseits in einer dem Relativitätspostulate nicht entsprechenden Weise angesetzt wird, also durch eine zufällige Kompensation zweier Widersprüche gegen das Relativitätspostulat. Indessen bedeutet diese Feststellung keinerlei Einwand gegen die molekulartheoretischen Hypothesen von Lorentz, sondern es wird nur klar, daß die Annahme der Kontraktion der Elektronen bei Bewegung in der Lorentzschen Theorie schon an einer früheren Stelle, als dieses durch Lorentz geschieht, eingeführt werden müßte.

In einem Anhange gehe ich noch auf die Stellung der klassischen Mechanik zum Relativitätspostulate ein. Eine leicht vorzunehmende Anpassung der Mechanik an das Relativitätspostulat würde für die beobachtbaren Erscheinungen kaum merkliche Differenzen ergeben, würde aber zu einem sehr überraschenden Erfolge führen: *Mit der Voranstellung des Relativitätspostulates schafft man sich genau das hinreichende Mittel, um hernach die vollständigen Gesetze der Mechanik allein aus dem Satze von der Erhaltung der Energie* (und Aussagen über die Formen der Energie) *zu entnehmen.*

§ 1.

Bezeichnungen.

Ein Bezugssystem x, y, z, t rechtwinkliger Koordinaten im Raume und der Zeit sei gegeben. Die Zeiteinheit soll in solcher Beziehung zur Längeneinheit gewählt sein, daß die Lichtgeschwindigkeit im leeren Raume 1 ist.

Obwohl ich an sich vorziehen würde, die von Lorentz gebrauchten Bezeichnungen nicht zu ändern, scheint es mir doch wichtig, gewisse Zusammengehörigkeiten durch eine andere Wahl der Zeichen von vornherein hervortreten zu lassen. Ich werde den Vektor

der *elektrischen Kraft* $\mathfrak{E}$, der *magnetischen Erregung* $\mathfrak{M}$, der *elektrischen Erregung* $\mathfrak{e}$, der *magnetischen Kraft* $\mathfrak{m}$

nennen, so daß also $\mathfrak{E}$, $\mathfrak{M}$, $\mathfrak{e}$, $\mathfrak{m}$ an die Stelle von $\mathfrak{E}$, $\mathfrak{B}$, $\mathfrak{D}$, $\mathfrak{H}$ bei Lorentz treten sollen.

Ich werde mich ferner des Gebrauchs komplexer Größen in einer Weise, wie dies bisher in physikalischen Untersuchungen noch nicht üblich war, bedienen, namentlich statt mit t mit der Verbindung it operieren, wobei i die imaginäre Einheit $\sqrt{-1}$ bedeute. Andererseits werden ganz wesentliche Umstände in Evidenz treten, indem ich eine Schreibweise mit Indizes benutzen werde, nämlich oft an Stelle von

$$x, y, z, it$$

$$x_1, x_2, x_3, x_4$$

setzen und hierauf einen allgemeinen Gebrauch der Indizes 1, 2, 3, 4 gründen werde. Dabei wird es sich, wie ich ausdrücklich hervorhebe, stets nur um eine *übersichtlichere Zusammenfassung rein reeller Beziehungen* handeln, und der Übergang zu reellen Gleichungen wird sich überall sofort vollziehen lassen, indem von den Zeichen mit Indizes solche mit *einem* Index 4 stets *rein imaginäre* Werte, solche mit *keinem* Index 4 oder mit *zwei* Indizes 4 stets *reelle* Werte bedeuten werden.

Ein einzelnes Wertsystem x, y, z, t bzw. x_1, x_2, x_3, x_4 soll ein *Raum-Zeitpunkt* heißen.

Ferner bezeichne w den Vektor *Geschwindigkeit der Materie*, ε die *Dielektrizitätskonstante*, μ die *magnetische Permeabilität*, σ die *Leitfähigkeit* der Materie, sämtlich als Funktionen von x, y, z, t (oder x_1, x_2, x_3, x_4) gedacht, weiter ϱ die *elektrische Raumdichte*, s einen Vektor „*elektrischer Strom*“, zu dessen Definition wir erst in der Folge (in § 7 und 8) kommen werden.

Erster Teil.

Betrachtung des Grenzfalles Äther.

§ 2.

Die Grundgleichungen für den Äther.

Die Lorentzsche Theorie führt die Gesetze der Elektrodynamik der ponderablen Körper durch atomistische Vorstellungen von der Elektrizität zurück auf einfachere Gesetze; an diese einfacheren Gesetze knüpfen wir hier ebenfalls an, indem wir fordern, daß sie den Grenzfall $\varepsilon = 1$, $\mu = 1$, $\sigma = 0$ der Gesetze für ponderable Körper bilden sollen. In diesem idealen Grenzfalle $\varepsilon = 1$, $\mu = 1$, $\sigma = 0$ soll $\mathfrak{E} = e$, $\mathfrak{M} = m$ sein und sollen an jedem Raum-Zeitpunkte x, y, z, t die Gleichungen bestehen:

$$(I) \quad \text{curl } m - \frac{\partial e}{\partial t} = \varrho w,$$

$$(II) \quad \text{div } e = \varrho,$$

$$(III) \quad \text{curl } e + \frac{\partial m}{\partial t} = 0,$$

$$(IV) \quad \text{div } m = 0.$$

Ich will nun schreiben x_1, x_2, x_3, x_4 für x, y, z, it ($i = \sqrt{-1}$), weiter

$$e_1, e_2, e_3, e_4$$

für

$$\varrho w_x, \varrho w_y, \varrho w_z, i\varrho,$$

d. s. die Komponenten des Konvektionsstromes $q\omega$ und die mit i multiplizierte Raumdichte der Elektrizität, ferner

$$\begin{aligned} & f_{23}, f_{31}, f_{12}, f_{14}, f_{24}, f_{34} \\ \text{für} & m_x, m_y, m_z, -ie_x, -ie_y, -ie_z, \end{aligned}$$

d. s. die Komponenten von m bzw. $-ie$ nach den Achsen, endlich noch allgemein bei zwei der Reihe 1, 2, 3, 4 entnommenen Indizes h, k

$$f_{kh} = -f_{hk},$$

also

$$\begin{aligned} f_{32} &= -f_{23}, f_{13} = -f_{31}, f_{21} = -f_{12}, \\ f_{41} &= -f_{14}, f_{42} = -f_{24}, f_{43} = -f_{34} \end{aligned}$$

festsetzen.

Als dann schreiben sich die drei in (I) zusammengefaßten Gleichungen und die mit i multiplizierte Gleichung (II):

$$\begin{aligned} & \frac{\partial f_{12}}{\partial x_2} + \frac{\partial f_{13}}{\partial x_3} + \frac{\partial f_{14}}{\partial x_4} = q_1, \\ (A) \quad & \frac{\partial f_{21}}{\partial x_1} + \frac{\partial f_{23}}{\partial x_3} + \frac{\partial f_{24}}{\partial x_4} = q_2, \\ & \frac{\partial f_{31}}{\partial x_1} + \frac{\partial f_{32}}{\partial x_2} + \frac{\partial f_{34}}{\partial x_4} = q_3, \\ & \frac{\partial f_{41}}{\partial x_1} + \frac{\partial f_{42}}{\partial x_2} + \frac{\partial f_{43}}{\partial x_3} = q_4. \end{aligned}$$

Andererseits verwandeln sich die drei in (III) zusammengefaßten Gleichungen, mit $-i$ multipliziert, und die Gleichung (IV), mit -1 multipliziert, in

$$\begin{aligned} & \frac{\partial f_{34}}{\partial x_2} + \frac{\partial f_{42}}{\partial x_3} + \frac{\partial f_{23}}{\partial x_4} = 0, \\ (B) \quad & \frac{\partial f_{43}}{\partial x_1} + \frac{\partial f_{14}}{\partial x_3} + \frac{\partial f_{31}}{\partial x_4} = 0, \\ & \frac{\partial f_{24}}{\partial x_1} + \frac{\partial f_{41}}{\partial x_2} + \frac{\partial f_{12}}{\partial x_4} = 0, \\ & \frac{\partial f_{32}}{\partial x_1} + \frac{\partial f_{13}}{\partial x_2} + \frac{\partial f_{21}}{\partial x_3} = 0. \end{aligned}$$

Man bemerkt bei dieser Schreibweise sofort die *vollkommene Symmetrie* des ersten wie des zweiten dieser Gleichungssysteme *in bezug auf die Permutationen der Indizes 1, 2, 3, 4.*

§ 3.

Das Theorem der Relativität von Lorentz.

Die Schreibweise der Gleichungen (I)–(IV) in der Symbolik des Vektorkalküls dient bekanntermaßen dazu, eine Invarianz (oder besser

Kovarianz) des Gleichungssystems (A) wie des Gleichungssystems (B) bei einer Drehung des Koordinatensystems um den Nullpunkt in Evidenz zu setzen. Nehmen wir z. B. eine Drehung um die z -Achse um einen festen Winkel φ vor unter Festhaltung der Vektoren $\mathbf{e}, \mathbf{m}, \mathbf{w}$ im Raume, führen also anstatt x_1, x_2, x_3, x_4 neue Variablen x_1', x_2', x_3', x_4' ein durch

$$x_1' = x_1 \cos \varphi + x_2 \sin \varphi, \quad x_2' = -x_1 \sin \varphi + x_2 \cos \varphi, \quad x_3' = x_3, \quad x_4' = x_4,$$

dazu neue Größen q_1', q_2', q_3', q_4' durch

$$q_1' = q_1 \cos \varphi + q_2 \sin \varphi, \quad q_2' = -q_1 \sin \varphi + q_2 \cos \varphi, \quad q_3' = q_3, \quad q_4' = q_4,$$

neue Größen $f_{12}', \dots, f_{34}'$ durch

$$\begin{aligned} f_{23}' &= f_{23} \cos \varphi + f_{31} \sin \varphi, \quad f_{31}' = -f_{23} \sin \varphi + f_{31} \cos \varphi, \quad f_{12}' = f_{12}, \\ f_{14}' &= f_{14} \cos \varphi + f_{24} \sin \varphi, \quad f_{24}' = -f_{14} \sin \varphi + f_{24} \cos \varphi, \quad f_{34}' = f_{34}, \\ f_{kh}' &= -f_{hk}' \quad (h, k = 1, 2, 3, 4), \end{aligned}$$

so wird notwendig aus (A) das genau entsprechende System (A'), aus (B) das genau entsprechende System (B') zwischen den neuen, mit Strichen versehenen Größen folgen.

Nun läßt sich auf Grund der Symmetrie des Systems (A) wie des Systems (B) in den Indizes 1, 2, 3, 4 sofort ohne jede Rechnung das von Lorentz gefundene Theorem der Relativität entnehmen.

Ich will unter $i\psi$ eine rein imaginäre Größe verstehen und die Substitution

$$(1) \quad x_1' = x_1, \quad x_2' = x_2, \quad x_3' = x_3 \cos i\psi + x_4 \sin i\psi, \quad x_4' = -x_3 \sin i\psi + x_4 \cos i\psi$$

betrachten. Mittels

$$(2) \quad -i \operatorname{tg} i\psi = \frac{e^\psi - e^{-\psi}}{e^\psi + e^{-\psi}} = q, \quad \psi = \frac{1}{2} \log \operatorname{nat} \frac{1+q}{1-q}$$

wird

$$\cos i\psi = \frac{1}{\sqrt{1-q^2}}, \quad \sin i\psi = \frac{iq}{\sqrt{1-q^2}},$$

wobei $-1 < q < 1$ ausfällt und $\sqrt{1-q^2}$ mit dem positiven Vorzeichen zu nehmen ist. Schreiben wir noch

$$(3) \quad x_1' = x', \quad x_2' = y', \quad x_3' = z', \quad x_4' = it',$$

so nimmt daher die Substitution (1) die Gestalt

$$(4) \quad x' = x, \quad y' = y, \quad z' = \frac{z - qt}{\sqrt{1-q^2}}, \quad t' = \frac{-qz + t}{\sqrt{1-q^2}}$$

mit lauter reellen Koeffizienten an.

Ersetzen wir nun in den oben bei der Drehung um die z -Achse genannten Gleichungen überall 1, 2, 3, 4 durch 3, 4, 1, 2, und gleichzeitig φ durch $i\psi$, so erkennen wir, daß wenn gleichzeitig mit dieser Substitution

(1) neue Größen q_1', q_2', q_3', q_4' durch

$$q_1' = q_1, \quad q_2' = q_2, \quad q_3' = q_3 \cos i\psi + q_4 \sin i\psi, \quad q_4' = -q_3 \sin i\psi + q_4 \cos i\psi,$$

neue Größen $f_{12}', \dots, f_{34}'$ durch

$$\begin{aligned} f_{41}' &= f_{41} \cos i\psi + f_{13} \sin i\psi, & f_{13}' &= -f_{41} \sin i\psi + f_{13} \cos i\psi, & f_{34}' &= f_{34}, \\ f_{32}' &= f_{32} \cos i\psi + f_{42} \sin i\psi, & f_{42}' &= -f_{32} \sin i\psi + f_{42} \cos i\psi, & f_{12}' &= f_{12}, \\ f_{hk}' &= -f_{hk}' & & & (h, k = 1, 2, 3, 4) \end{aligned}$$

eingeführt werden, alsdann ebenfalls das System (A) in das genau entsprechende System (A'), das System (B) in das genau entsprechende System (B') zwischen den neuen, mit Strichen versehenen Größen übergehen wird.

Alle diese Gleichungen lassen sich sofort in rein reelle Gestalt umschreiben und man kann das letzte Ergebnis so formulieren:

Wird die reelle Transformation (4) vorgenommen und werden hernach x', y', z', t' als ein Bezugssystem für Raum und Zeit angesprochen, werden zugleich

$$(5) \quad \varrho' = \varrho \left(\frac{-q w_z + 1}{\sqrt{1 - q^2}} \right), \quad \varrho' w_z' = \varrho \left(\frac{w_z - q}{\sqrt{1 - q^2}} \right), \quad \varrho' w_x' = \varrho w_x, \quad \varrho' w_y' = \varrho w_y,$$

ferner

$$(6) \quad e_x' = \frac{e_x - q m_y}{\sqrt{1 - q^2}}, \quad m_y' = \frac{-q e_x + m_y}{\sqrt{1 - q^2}}, \quad e_z' = e_z$$

und

$$(7) \quad m_x' = \frac{m_x + q e_y}{\sqrt{1 - q^2}}, \quad e_y' = \frac{q m_x + e_y}{\sqrt{1 - q^2}}, \quad m_z' = m_z$$

eingeführt*), so kommen hernach für die Vektoren w', e', m' mit den Komponenten $w_x', w_y', w_z'; e_x', e_y', e_z'; m_x', m_y', m_z'$ in dem neuen Koordinatensystem x', y', z' und dazu die Größe ϱ' genau die zu (I)—(IV) analogen Gleichungen (I')—(IV') zustande, und zwar geht für sich das System (I), (II) in (I'), (II'), das System (III), (IV) in (III'), (IV') über.

Wir bemerken, daß hier $e_x - q m_y, e_y + q m_x, e_z$ die Komponenten des Vektors $e + [vm]$ sind, wenn v einen Vektor in Richtung der positiven z -Achse vom Betrage $|v| = q$ und $[vm]$ das vektorielle Produkt der Vektoren v und m bedeutet. Analog sind dann $m_x + q e_y, m_y - q e_x, m_z$ die Komponenten des Vektors $m - [ve]$.

Die Gleichungen (6) und (7), wie sie paarweise *unter* einander stehen, können durch eine andere Verwendung imaginärer Größen in

*) Die Gleichungen (5) stehen hier in anderer Folge, die Gleichungen (6) und (7) aber in der nämlichen Folge wie die zuvor genannten Gleichungen, die auf sie hinauskommen.

$$\begin{aligned}e'_x + im'_x &= (e_x + im_x) \cos i\psi + (e_y + im_y) \sin i\psi, \\e'_y + im'_y &= -(e_x + im_x) \sin i\psi + (e_y + im_y) \cos i\psi, \\e'_z + im'_z &= e_z + im_z\end{aligned}$$

zusammengefaßt werden, und wir merken noch an, daß, wenn φ irgendeinen reellen Winkel bedeutet, aus diesen letzten Beziehungen ferner die Kombinationen

$$(8) \quad \begin{aligned}&(e'_x + im'_x) \cos \varphi + (e'_y + im'_y) \sin \varphi \\&= (e_x + im_x) \cos (\varphi + i\psi) + (e_y + im_y) \sin (\varphi + i\psi),\end{aligned}$$

$$(9) \quad \begin{aligned}&-(e'_x + im'_x) \sin \varphi + (e'_y + im'_y) \cos \varphi \\&= -(e_x + im_x) \sin (\varphi + i\psi) + (e_y + im_y) \cos (\varphi + i\psi)\end{aligned}$$

hervorgehen.

§ 4.

Spezielle Lorentz-Transformationen.

Die Rolle, welche die z -Richtung in der Transformation (4) spielt, kann leicht auf eine beliebige Richtung übertragen werden, indem sowohl das Achsensystem der x, y, z wie das der x', y', z' , jedes einer und der nämlichen Drehung in bezug auf sich unterworfen wird. Wir kommen damit zu einem allgemeineren Satze.

Es sei $\mathbf{v}$ mit den Komponenten v_x, v_y, v_z ein gegebener Vektor mit einem solchen von Null verschiedenen Betrage $|\mathbf{v}| = q$, der kleiner als 1 ist, von irgendeiner Richtung. Wir verstehen allgemein unter $\bar{\mathbf{v}}$ eine beliebige auf $\mathbf{v}$ senkrechte Richtung und bezeichnen ferner die Komponente eines Vektors $\mathbf{r}$ nach der Richtung $\mathbf{v}$ oder einer Richtung $\bar{\mathbf{v}}$ mit r_v bzw. $r_{\bar{v}}$.

Anstatt x, y, z, t sollen nun neue Größen x', y', z', t' in folgender Weise eingeführt werden. Wird kurz $\mathbf{r}$ für den Vektor mit den Komponenten x, y, z im ersten, ferner $\mathbf{r}'$ für den Vektor mit den Komponenten x', y', z' im zweiten Bezugssystem geschrieben, so soll sein für die Richtung von $\mathbf{v}$:

$$(10) \quad r'_v = \frac{r_v - qt}{\sqrt{1 - q^2}},$$

für jede auf $\mathbf{v}$ senkrechte Richtung $\bar{\mathbf{v}}$:

$$(11) \quad r'_{\bar{v}} = r_{\bar{v}},$$

und ferner:

$$(12) \quad t' = \frac{-qr_v + t}{\sqrt{1 - q^2}}.$$

Die Bezeichnungen $\mathbf{r}'_v$ und $\mathbf{r}'_{\bar{v}}$ hier sind in dem Sinne zu verstehen, daß der Richtung v und jeder zu v senkrechten Richtung $\bar{v}$ in x, y, z immer die Richtung mit den nämlichen Richtungskosinus in x', y', z' zugeordnet wird.

Eine Transformation, wie sie durch (10), (11), (12) mit der Bedingung $0 < q < 1$ dargestellt wird, will ich eine *spezielle Lorentz-Transformation* nennen, und soll v der *Vektor*, die Richtung von v die *Achse*, der Betrag von v das *Moment* dieser speziellen Lorentz-Transformation heißen.

Werden weiter ϱ' und die Vektoren w', e', m' in x', y', z' dadurch definiert, daß

$$(13) \quad \varrho' = \frac{e(-qv_v + 1)}{\sqrt{1 - q^2}},$$

$$(14) \quad \varrho' w'_v = \frac{e w_v - e q}{\sqrt{1 - q^2}}, \quad \varrho' w'_v = \varrho w_{\bar{v}},$$

ferner*)

$$(15) \quad (e' + i m')_{\bar{v}} = \frac{(e + i m - i[v, e + i m]_{\bar{v}})}{\sqrt{1 - q^2}},$$

$$(e' + i m')_v = (e + i m - i[v, e + i m])_v$$

ist, so folgt der Satz, daß das Gleichungssystem (I), (II) und (III), (IV) jedesmal in das genau entsprechende System zwischen den mit Strichen versehenen Größen übergeht.

Die Auflösung der Gleichungen (10), (11), (12) führt auf:

$$(16) \quad \mathbf{r}_v = \frac{\mathbf{r}'_v + q \mathbf{t}'}{\sqrt{1 - q^2}}, \quad \mathbf{r}_{\bar{v}} = \mathbf{r}'_{\bar{v}}, \quad t = \frac{q \mathbf{r}'_v + \mathbf{t}'}{\sqrt{1 - q^2}}. \quad -$$

Wir schließen nun eine in der Folge sehr wichtige Bemerkung über die Beziehung der Vektoren w und w' an. Es möge wieder die schon mehrfach gebrauchte Bezeichnung mit den Indizes 1, 2, 3, 4 herangezogen werden, so daß wir x'_1, x'_2, x'_3, x'_4 für x', y', z', it' und $\varrho'_1, \varrho'_2, \varrho'_3, \varrho'_4$ für $\varrho' w'_x, \varrho' w'_y, \varrho' w'_z, i \varrho'$ setzen. Wie eine Drehung um die z -Achse, so ist offenbar auch die Transformation (4) und allgemeiner die Transformation (10), (11), (12) eine solche *lineare Transformation* von der Determinante $+1$, wodurch

$$(17) \quad x_1^2 + x_2^2 + x_3^2 + x_4^2, \text{ d. i. } x^2 + y^2 + z^2 - t^2$$

in

$$x_1'^2 + x_2'^2 + x_3'^2 + x_4'^2, \text{ d. i. } x'^2 + y'^2 + z'^2 - t'^2$$

übergeht.

*) Die runden Klammern sollen nur die Ausdrücke zusammenfassen, welche der Index betrifft, und $[v, e + i m]$ soll das vektorielle Produkt von v und $e + i m$ bedeuten.

Es wird daher auf Grund der Ausdrücke (13), (14) auch

$$-(\varrho_1^2 + \varrho_2^2 + \varrho_3^2 + \varrho_4^2) = \varrho^2 (1 - w_x^2 - w_y^2 - w_z^2) = \varrho^2 (1 - w^2)$$

in $\varrho'^2 (1 - w'^2)$ übergehen, oder mit andern Worten

$$(18) \quad \varrho \sqrt{1 - w^2},$$

wobei die Quadratwurzel positiv genommen sei, eine *Invariante* bei Lorentz-Transformationen sein.

Indem wir $\varrho_1, \varrho_2, \varrho_3, \varrho_4$ durch diese GröÙe dividieren, entstehen die 4 Werte

$$w_1 = \frac{w_x}{\sqrt{1 - w^2}}, \quad w_2 = \frac{w_y}{\sqrt{1 - w^2}}, \quad w_3 = \frac{w_z}{\sqrt{1 - w^2}}, \quad w_4 = \frac{i}{\sqrt{1 - w^2}},$$

zwischen welchen die Beziehung

$$(19) \quad w_1^2 + w_2^2 + w_3^2 + w_4^2 = -1$$

besteht. Offenbar sind diese 4 Werte eindeutig durch den Vektor w bestimmt, und umgekehrt folgt aus irgend 4 Werten w_1, w_2, w_3, w_4 , wobei w_1, w_2, w_3 reell, $-iw_4$ reell und positiv ist und die Bedingung (19) statthat, rückwärts gemäß diesen Gleichungen eindeutig ein Vektor w von einem Betrage < 1 .

Die Bedeutung von w_1, w_2, w_3, w_4 hier ist, daß sie die Verhältnisse von dx_1, dx_2, dx_3, dx_4 zu

$$(20) \quad \sqrt{-(dx_1^2 + dx_2^2 + dx_3^2 + dx_4^2)} = dt \sqrt{1 - w^2}$$

für die im Raum-Zeitpunkte x_1, x_2, x_3, x_4 befindliche Materie beim Übergang zu zeitlich benachbarten Zuständen derselben Stelle der Materie sind. Nun übertragen sich die Gleichungen (10), (11), (12) sofort auf die zusammengehörigen Differentiale dx, dy, dz, dt und dx', dy', dz', dt' und insbesondere wird daher für sie

$$-(dx_1^2 + dx_2^2 + dx_3^2 + dx_4^2) = -(dx_1'^2 + dx_2'^2 + dx_3'^2 + dx_4'^2)$$

sein. Nach Ausführung der Lorentz-Transformation ist im neuen Bezugssystem als Geschwindigkeit der Materie im nämlichen Raum-Zeitpunkte x', y', z', t' der Vektor w' mit den Verhältnissen $\frac{dx'}{dt'}, \frac{dy'}{dt'}, \frac{dz'}{dt'}$ als Komponenten auszulegen.

Nunmehr ist ersichtlich, daß das Wertsystem

$$x_1 = w_1, \quad x_2 = w_2, \quad x_3 = w_3, \quad x_4 = w_4$$

vermöge der Lorentz-Transformation (10), (11), (12) eben in dasjenige neue Wertsystem

$$x_1' = w_1', \quad x_2' = w_2', \quad x_3' = w_3', \quad x_4' = w_4'$$

übergeht, das für die Geschwindigkeit w' nach der Transformation genau die Bedeutung hat wie das erstere Wertsystem für die Geschwindigkeit vor der Transformation.

Ist insbesondere der Vektor v der speziellen Lorentz-Transformation gleich dem Geschwindigkeitsvektor w der Materie im Raum-Zeitpunkte x_1, x_2, x_3, x_4 , so folgt aus (10), (11), (12):

$$w_1' = 0, \quad w_2' = 0, \quad w_3' = 0, \quad w_4' = i.$$

Unter diesen Umständen erhält also der betreffende Raum-Zeitpunkt nach der Transformation die Geschwindigkeit $w' = 0$, er wird, wie wir uns ausdrücken können, *auf Ruhe transformiert*. Wir können danach die Invariante $\varrho\sqrt{1-w^2}$ passend als *Ruh-Dichte* der Elektrizität bezeichnen.

§ 5.

Raum-Zeit-Vektoren I^{ter} und II^{ter} Art.

Indem wir das Hauptergebnis bezüglich der speziellen Lorentz-Transformationen mit der Tatsache zusammennehmen, daß das System (A) wie das System (B) jedenfalls bei einer Drehung des räumlichen Bezugssystems um den Nullpunkt kovariant ist, erhalten wir das allgemeine *Theorem der Relativität*. Um es leicht verständlich zu formulieren, dürfte es zweckmäßig sein, zuvor eine Reihe von abkürzenden Ausdrücken festzulegen, während ich andererseits daran festhalten will, komplexe Größen zu verwenden, um bestimmte Symmetrien in Evidenz zu setzen.

Eine lineare homogene Transformation

$$(21) \quad \begin{aligned} x_1 &= \alpha_{11}x_1' + \alpha_{12}x_2' + \alpha_{13}x_3' + \alpha_{14}x_4', \\ x_2 &= \alpha_{21}x_1' + \alpha_{22}x_2' + \alpha_{23}x_3' + \alpha_{24}x_4', \\ x_3 &= \alpha_{31}x_1' + \alpha_{32}x_2' + \alpha_{33}x_3' + \alpha_{34}x_4', \\ x_4 &= \alpha_{41}x_1' + \alpha_{42}x_2' + \alpha_{43}x_3' + \alpha_{44}x_4' \end{aligned}$$

von der Determinante $+1$, in welcher alle Koeffizienten ohne einen Index 4 reell, dagegen $\alpha_{14}, \alpha_{24}, \alpha_{34}$ sowie $\alpha_{41}, \alpha_{42}, \alpha_{43}$ rein imaginär (ev. Null), endlich α_{44} wieder reell und speziell > 0 ist und durch welche

$$x_1^2 + x_2^2 + x_3^2 + x_4^2 \text{ in } x_1'^2 + x_2'^2 + x_3'^2 + x_4'^2$$

übergeht, will ich allgemein eine *Lorentz-Transformation* nennen.

Wird

$$x_1' = x', \quad x_2' = y', \quad x_3' = z', \quad x_4' = it'$$

gesetzt, so entsteht daraus sofort eine homogene lineare Transformation von x, y, z, t in x', y', z', t' mit lauter reellen Koeffizienten, wobei das Aggregat

$$-x^2 - y^2 - z^2 + t^2 \text{ in } -x'^2 - y'^2 - z'^2 + t'^2$$

übergeht und einem jeden solchen Wertesystem x, y, z, t mit *positivem* t , wofür dieses Aggregat > 0 ausfällt, stets auch ein *positives* t' entspricht; letzteres ist aus der Kontinuität des Aggregats in x, y, z, t leicht ersichtlich.

Die letzte Vertikalreihe des Koeffizientensystems von (21) hat die Bedingung

$$(22) \quad \alpha_{14}^2 + \alpha_{24}^2 + \alpha_{34}^2 + \alpha_{44}^2 = 1$$

zu erfüllen.

Sind $\alpha_{14} = 0$, $\alpha_{24} = 0$, $\alpha_{34} = 0$, so ist $\alpha_{44} = 1$ und die Lorentz-Transformation reduziert sich auf eine bloße Drehung des räumlichen Koordinatensystems um den Nullpunkt.

Sind α_{14} , α_{24} , α_{34} nicht sämtlich Null und setzt man

$$\alpha_{14} : \alpha_{24} : \alpha_{34} : \alpha_{44} = v_x : v_y : v_z : i,$$

so folgt aus (22) der Betrag

$$q = \sqrt{v_x^2 + v_y^2 + v_z^2} < 1.$$

Andererseits kann man zu jedem Wertesystem α_{14} , α_{24} , α_{34} , α_{44} , das in dieser Weise mit reellen v_x , v_y , v_z die Bedingung (22) erfüllt, die *spezielle* Lorentz-Transformation (16) mit α_{14} , α_{24} , α_{34} , α_{44} als letzter Vertikalreihe konstruieren und jede Lorentz-Transformation mit der nämlichen letzten Vertikalreihe der Koeffizienten kann alsdann zusammengesetzt werden aus dieser speziellen Lorentz-Transformation und einer sich daran anschließenden Drehung des räumlichen Koordinatensystems um den Nullpunkt.

Die Gesamtheit aller Lorentz-Transformationen bildet eine *Gruppe*.

Unter einem *Raum-Zeit-Vektor* I. Art soll verstanden werden ein beliebiges System von vier Größen q_1, q_2, q_3, q_4 mit der Vorschrift, bei jeder Lorentz-Transformation (21) es durch dasjenige System q'_1, q'_2, q'_3, q'_4 zu ersetzen, das aus (21) für die Werte x'_1, x'_2, x'_3, x'_4 hervorgeht, wenn für x_1, x_2, x_3, x_4 die Werte q_1, q_2, q_3, q_4 genommen werden.

Verwenden wir neben dem variablen Raum-Zeit-Vektor I. Art x_1, x_2, x_3, x_4 einen zweiten solchen variablen Raum-Zeit-Vektor I. Art u_1, u_2, u_3, u_4 und fassen die bilineare Verbindung

$$(23) \quad \begin{aligned} & f_{23}(x_2 u_3 - x_3 u_2) + f_{31}(x_3 u_1 - x_1 u_3) + f_{12}(x_1 u_2 - x_2 u_1) \\ & + f_{14}(x_1 u_4 - x_4 u_1) + f_{24}(x_2 u_4 - x_4 u_2) + f_{34}(x_3 u_4 - x_4 u_3) \end{aligned}$$

mit sechs Koeffizienten $f_{23}, \dots, f_{34}$ auf. Wir bemerken, daß diese einerseits sich in vektorieller Schreibweise aus den vier Vektoren

$$x_1, x_2, x_3; \quad u_1, u_2, u_3; \quad f_{23}, f_{31}, f_{12}; \quad f_{14}, f_{24}, f_{34}$$

und den Konstanten x_4 und u_4 aufbauen läßt, andererseits symmetrisch in den Indizes 1, 2, 3, 4 ist. Indem wir x_1, x_2, x_3, x_4 und u_1, u_2, u_3, u_4

gleichzeitig gemäß der Lorentz-Transformation (21) substituieren, geht (23) in eine Verbindung

$$(24) \quad \begin{aligned} & f'_{23}(x'_2 u'_3 - x'_3 u'_2) + f'_{31}(x'_3 u'_1 - x'_1 u'_3) + f'_{12}(x'_1 u'_2 - x'_2 u'_1) \\ & + f'_{14}(x'_1 u'_4 - x'_4 u'_1) + f'_{24}(x'_2 u'_4 - x'_4 u'_2) + f'_{34}(x'_3 u'_4 - x'_4 u'_3) \end{aligned}$$

mit gewissen allein von den sechs Größen $f_{23}, \dots, f_{34}$ und den sechzehn Koeffizienten $\alpha_{11}, \alpha_{12}, \dots, \alpha_{44}$ abhängenden sechs Koeffizienten $f'_{23}, \dots, f'_{34}$ über.

Einen *Raum-Zeit-Vektor II. Art* definieren wir als ein System von sechs Größen $f_{23}, f_{31}, f_{12}, f_{14}, f_{24}, f_{34}$ mit der Vorschrift, es bei jeder Lorentz-Transformation durch dasjenige neue System $f'_{23}, f'_{31}, f'_{12}, f'_{14}, f'_{24}, f'_{34}$ zu ersetzen, das dem eben erörterten Zusammenhange der Form (23) mit der Form (24) entspricht.

Das allgemeine Theorem der Relativität betreffend die Gleichungen (I)–(IV), die „Grundgleichungen für den Äther“, spreche ich nunmehr folgendermaßen aus.

Werden x, y, z, it (Raumkoordinaten und Zeit $\propto i$) einer beliebigen Lorentz-Transformation unterworfen und gleichzeitig $\varrho w_x, \varrho w_y, \varrho w_z, i\varrho$ (Konvektionsstrom und Ladungsdichte $\propto i$) als *Raum-Zeit-Vektor I. Art*, ferner $m_x, m_y, m_z, -ie_x, -ie_y, -ie_z$ (magnetische Kraft und elektrische Erregung $\propto -i$) als *Raum-Zeit-Vektor II. Art* transformiert, so geht das System der Gleichungen (I), (II) und das System der Gleichungen (III), (IV) je in das System der entsprechend lautenden Beziehungen zwischen den entsprechenden neu eingeführten Größen über.

Kürzer mag diese Tatsache auch mit den Worten angedeutet werden: Das System der Gleichungen (I), (II) wie das System der Gleichungen (III), (IV) ist kovariant bei jeder Lorentz-Transformation, wobei $\varrho w, i\varrho$ als *Raum-Zeit-Vektor I. Art*, $m, -ie$ als *Raum-Zeit-Vektor II. Art* zu transformieren ist. Oder noch prägnanter:

$\varrho w, i\varrho$ ist ein *Raum-Zeit-Vektor I. Art*, $m, -ie$ ist ein *Raum-Zeit-Vektor II. Art*. —

Ich füge noch einige Bemerkungen hier an, um die Vorstellung eines *Raum-Zeit-Vektors II. Art* zu erleichtern. *Invarianten* für einen solchen Vektor $m, -ie$ bei der Gruppe der Lorentz-Transformationen sind offenbar

$$(25) \quad m^2 - e^2 = f_{23}^2 + f_{31}^2 + f_{12}^2 + f_{14}^2 + f_{24}^2 + f_{34}^2,$$

$$(26) \quad m e = i(f_{23} f_{14} + f_{31} f_{24} + f_{12} f_{34}).$$

Ein *Raum-Zeit-Vektor II. Art* $m, -ie$, (wobei m und e reelle Raum-Vektoren sind), mag *singulär* heißen, wenn das skalare Quadrat $(m - ie)^2 = 0$, d. h. $m^2 - e^2 = 0$ und zugleich $(me) = 0$ ist, d. h. *die Vektoren m und e gleichen Betrag haben und zudem senkrecht aufeinander stehen*. Wenn solches der Fall ist, bleiben diese zwei Eigenschaften für den *Raum-Zeit-Vektor II. Art* bei jeder Lorentz-Transformation erhalten.

Ist der Raum-Zeit-Vektor II. Art m , — *ie nicht singular*, so drehen wir zunächst das räumliche Koordinatensystem so, daß das Vektorprodukt $[m\epsilon]$ in die z -Achse fällt, daß $m_z = 0$, $\epsilon_z = 0$ ist. Dann ist

$$(m_x - i\epsilon_x)^2 + (m_y - i\epsilon_y)^2 = 0,$$

also $\frac{\epsilon_y + im_y}{\epsilon_x + im_x}$ verschieden von $\pm i$ und wir können daher ein komplexes Argument $\varphi + i\psi$ derart bestimmen, daß

$$\operatorname{tg}(\varphi + i\psi) = \frac{\epsilon_y + im_y}{\epsilon_x + im_x}$$

ist. Alsdann wird mit Rücksicht auf die Gleichung (9) durch die zu ψ gehörige Transformation (1) und eine nachherige Drehung um die z -Achse durch den Winkel φ eine Lorentz-Transformation bewirkt, nach der auch noch $m_y = 0$, $\epsilon_y = 0$ werden, also nunmehr m und ϵ beide in die neue x -Linie fallen; dabei sind durch die Invarianten $m^2 - \epsilon^2$ und $(m\epsilon)$ die schließlichen Größen dieser Vektoren und ob sie von gleicher oder entgegengesetzter Richtung werden oder einer Null wird, von vornherein fixiert.

§ 6.

Begriff der Zeit.

Durch die Lorentz-Transformationen werden gewisse Abänderungen des Zeitparameters zugelassen. Infolgedessen ist es nicht mehr statthaft, von der *Gleichzeitigkeit* zweier Ereignisse an sich zu sprechen. Die Verwendung dieses Begriffs setzt vielmehr voraus, daß die Freiheit der 6 Parameter, die zur Angabe eines Bezugssystems für Raum und Zeit offen steht, bereits in gewisser Weise auf eine Freiheit von nur 3 Parametern eingeschränkt ist. Nur weil wir gewohnt sind, diese Einschränkung stark approximativ eindeutig zu treffen, halten wir den Begriff der Gleichzeitigkeit zweier Ereignisse als an sich existierend.*) In Wahrheit aber sollen folgende Umstände zutreffen.

Ein Bezugssystem x, y, z, t für Raum-Zeitpunkte (Ereignisse) sei irgendwie bekannt. Wird ein Raumpunkt $A(x_0, y_0, z_0)$ zur Zeit $t_0 = 0$ mit einem anderen Raumpunkte $P(x, y, z)$ zu einer anderen Zeit t verglichen und ist die Zeitdifferenz $t - t_0$ (es sei etwa $t > t_0$) *kleiner* als die Länge AP , d. i. die Zeit, die das Licht zur Fortpflanzung von A nach P braucht, und ist q der Quotient $\frac{t - t_0}{AP} < 1$, so können wir durch die spezielle Lorentz-Transformation, die AP als Achse und q als Moment

*) Ungefähr wie Wesen, gebannt an eine enge Umgebung eines Punktes auf einer Kugeloberfläche, darauf verfallen könnten, die Kugel sei ein geometrisches Gebilde, an welchem ein Durchmesser an sich ausgezeichnet ist.

hat, einen neuen Zeitparameter t' einführen, der (s. Gleichung (12) in § 4) für beide Raum-Zeitpunkte A, t_0 und P, t den gleichen Wert $t' = 0$ erlangt; es lassen sich also diese zwei Ereignisse auch als gleichzeitig auffassen.

Nehmen wir weiter zu einer und derselben Zeit $t_0 = 0$ zwei verschiedene Raumpunkte A, B oder drei Raumpunkte A, B, C , die nicht in einer Geraden liegen, und vergleichen damit einen Raumpunkt P außerhalb der Geraden AB oder der Ebene ABC zu einer anderen Zeit t und ist die Zeitdifferenz $t - t_0$ (es sei etwa $t > t_0$) *kleiner* als die Zeit, die das Licht zur Fortpflanzung von der Geraden AB oder der Ebene ABC nach P braucht, und q der Quotient aus der ersteren und der letzteren Zeit, so erscheinen nach Anwendung der speziellen Lorentz-Transformation, die als Achse das Lot auf AB , bzw. ABC durch P und als Moment q hat, alle drei (beziehungsweise vier) Ereignisse $A, t_0; B, t_0; (C, t_0)$ und P, t als gleichzeitig.

Werden jedoch vier Raumpunkte, die nicht in einer Ebene liegen, zu einer und derselben Zeit t_0 aufgefaßt, so ist es nicht mehr möglich, durch eine Lorentz-Transformation eine Abänderung des Zeitparameters vorzunehmen, ohne daß der Charakter der Gleichzeitigkeit dieser vier Raum-Zeitpunkte verloren geht.

Dem Mathematiker, der an Betrachtungen über mehrdimensionale Mannigfaltigkeiten und andererseits an die Begriffsbildungen der sogenannten nicht-Euklidischen Geometrie gewöhnt ist, kann es keine wesentliche Schwierigkeit bereiten, den Begriff der Zeit an die Verwendung der Lorentz-Transformationen zu adaptieren. Dem Bedürfnisse, sich das Wesen dieser Transformationen physikalisch näher zu bringen, kommt der in der Einleitung zitierte Aufsatz von A. Einstein entgegen.

Zweiter Teil.

Die elektromagnetischen Vorgänge.

§ 7.

Die Grundgleichungen für ruhende Körper.

Nach diesen vorbereitenden Ausführungen, die wir des etwas geringeren mathematischen Apparates wegen an dem idealen Grenzfalle $\epsilon = 1$, $\mu = 1$, $\sigma = 0$ entwickelten, wenden wir uns jetzt zu den Gesetzen für die elektromagnetischen Vorgänge in der Materie. Wir suchen diejenigen Beziehungen, die es — unter Voraussetzung geeigneter Grenzdaten — ermöglichen, an jedem Orte und zu jeder Zeit, also als Funktionen von x, y, z, t zu finden: die Vektoren der elektrischen Kraft $\mathfrak{E}$, der magne-

tischen Erregung $\mathfrak{M}$, der elektrischen Erregung $\mathfrak{e}$, der magnetischen Kraft $\mathfrak{m}$, die elektrische Raumdichte ϱ , den Vektor „elektrischer Strom $\mathfrak{s}$ “, (dessen Beziehung zum Leitungsstrom hernach durch die Art des Auftretens der Leitfähigkeit zu erkennen sein wird), endlich den Vektor $\mathfrak{w}$, die Geschwindigkeit der Materie.

Die fraglichen Beziehungen scheiden sich in zwei Klassen,

erstens diejenigen Gleichungen, die, wenn der Vektor $\mathfrak{w}$ als Funktion von x, y, z, t gegeben, also die Bewegung der Materie bekannt ist, zur Kenntnis aller anderen eben genannten Größen als Funktionen von x, y, z, t hinführen, — diese erste Klasse speziell will ich die *Grundgleichungen* nennen, —

zweitens die Ausdrücke für die *ponderomotorischen Kräfte*, die durch Heranziehen der Gesetze der Mechanik weiter Aufschluß über den Vektor $\mathfrak{w}$ als Funktion von x, y, z, t bringen.

Für den Fall *ruhender Körper*, d. i. wenn $\mathfrak{w}(x, y, z, t) = 0$ gegeben ist, kommen die Theorien von Maxwell (Heaviside, Hertz) und von Lorentz zu den nämlichen Grundgleichungen. Es sind dies

1) die *Differentialgleichungen*, die noch keine auf die Materie bezüglichen Konstanten enthalten:

$$(I) \quad \text{curl } \mathfrak{m} - \frac{\partial \mathfrak{e}}{\partial t} = \mathfrak{s},$$

$$(II) \quad \text{div } \mathfrak{e} = \varrho,$$

$$(III) \quad \text{curl } \mathfrak{E} + \frac{\partial \mathfrak{M}}{\partial t} = 0,$$

$$(IV) \quad \text{div } \mathfrak{M} = 0;$$

2) weitere Beziehungen, die den Einfluß der vorhandenen Materie charakterisieren; sie werden in dem wichtigsten Falle, auf den wir uns hier beschränken, für isotrope Körper, angesetzt in der Gestalt

$$(V) \quad \mathfrak{e} = \varepsilon \mathfrak{E}, \quad \mathfrak{M} = \mu \mathfrak{m}, \quad \mathfrak{s} = \sigma \mathfrak{E},$$

wobei ε die Dielektrizitätskonstante, μ die magnetische Permeabilität, σ die Leitfähigkeit der Materie als Funktionen von x, y, z und t bekannt zu denken sind. $\mathfrak{s}$ ist hier als *Leitungsstrom* anzusprechen.

Ich lasse nun an diesen Gleichungen wieder durch eine veränderte Schreibweise eine noch versteckte Symmetrie hervortreten. Ich setze wie in den vorangeschickten Ausführungen

$$x_1 = x, \quad x_2 = y, \quad x_3 = z, \quad x_4 = it$$

und schreibe

$$s_1, s_2, s_3, s_4$$

für

$$\mathfrak{s}_x, \mathfrak{s}_y, \mathfrak{s}_z, i\varrho,$$

ferner

$$f_{23}, f_{31}, f_{12}, f_{14}, f_{24}, f_{34}$$

für

$$m_x, m_y, m_z, -ie_x, -ie_y, -ie_z$$

und noch

$$F_{23}, F_{31}, F_{12}, F_{14}, F_{24}, F_{34}$$

für

$$\mathcal{M}_x, \mathcal{M}_y, \mathcal{M}_z, -i\mathcal{E}_x, -i\mathcal{E}_y, -i\mathcal{E}_z;$$

endlich soll für andere Paare von ungleichen, der Reihe 1, 2, 3, 4 entnommenen Indizes h, k stets

$$f_{kh} = -f_{hk}, \quad F_{kh} = -F_{hk}$$

gelten. (Die Buchstaben f, F sollen an das Wort Feld, s an Strom erinnern.)

Dann schreiben sich die Gleichungen (I), (II) um in

$$\begin{aligned} & \frac{\partial f_{12}}{\partial x_2} + \frac{\partial f_{13}}{\partial x_3} + \frac{\partial f_{14}}{\partial x_4} = s_1, \\ (A) \quad & \frac{\partial f_{21}}{\partial x_1} + \frac{\partial f_{23}}{\partial x_3} + \frac{\partial f_{24}}{\partial x_4} = s_2, \\ & \frac{\partial f_{31}}{\partial x_1} + \frac{\partial f_{32}}{\partial x_2} + \frac{\partial f_{34}}{\partial x_4} = s_3, \\ & \frac{\partial f_{41}}{\partial x_1} + \frac{\partial f_{42}}{\partial x_2} + \frac{\partial f_{43}}{\partial x_3} = s_4, \end{aligned}$$

und die Gleichungen (III), (IV) schreiben sich um in

$$\begin{aligned} & \frac{\partial F_{34}}{\partial x_2} + \frac{\partial F_{42}}{\partial x_3} + \frac{\partial F_{23}}{\partial x_4} = 0, \\ (B) \quad & \frac{\partial F_{43}}{\partial x_1} + \frac{\partial F_{14}}{\partial x_3} + \frac{\partial F_{31}}{\partial x_4} = 0, \\ & \frac{\partial F_{24}}{\partial x_1} + \frac{\partial F_{41}}{\partial x_2} + \frac{\partial F_{12}}{\partial x_4} = 0, \\ & \frac{\partial F_{32}}{\partial x_1} + \frac{\partial F_{13}}{\partial x_2} + \frac{\partial F_{21}}{\partial x_3} = 0. \end{aligned}$$

§ 8.

Die Grundgleichungen für bewegte Körper.

Nunmehr wird es uns gelingen, die Grundgleichungen für beliebig bewegte Körper in eindeutiger Weise festzustellen, ausschließlich mittels folgender drei Axiome:

Das erste Axiom soll sein:

Wenn eine einzelne Stelle der Materie in einem Momente ruht, also der Vektor m für ein System x, y, z, t Null ist, — die Umgebung mag in irgendwelcher Bewegung begriffen sein —, so sollen für den Raum-

Zeitpunkt x, y, z, t zwischen ϱ , den Vektoren $\mathfrak{s}, \mathfrak{e}, \mathfrak{m}, \mathfrak{E}, \mathfrak{M}$ und deren Ableitungen nach x, y, z, t genau die Beziehungen (A), (B), (V) statt haben, die zu gelten hätten, falls alle Materie ruhte.

Das zweite Axiom soll sein:

Jede Geschwindigkeit der Materie ist < 1 , kleiner als die Fortpflanzungsgeschwindigkeit des Lichtes im leeren Raume.

Das dritte Axiom soll sein:

Die Grundgleichungen sind von solcher Art, daß wenn x, y, z, t irgendeiner Lorentz-Transformation unterworfen und dabei einerseits $\mathfrak{m}, -i\mathfrak{e}$, andererseits $\mathfrak{M}, -i\mathfrak{E}$ je als Raum-Zeit-Vektor II. Art, $\mathfrak{s}, i\varrho$ als Raum-Zeit-Vektor I. Art transformiert werden, die Gleichungen dadurch in die genau entsprechend lautenden Gleichungen zwischen den transformierten Größen übergehen.

Dieses dritte Axiom deute ich auch kurz mit den Worten an:

$\mathfrak{m}, -i\mathfrak{e}$ und $\mathfrak{M}, -i\mathfrak{E}$ sind je ein Raum-Zeit-Vektor II. Art, $\mathfrak{s}, i\varrho$ ein Raum-Zeit-Vektor I. Art, und dieses Axiom nenne ich das *Prinzip der Relativität*.

Diese drei Axiome führen uns in der Tat von den vorhin genannten Grundgleichungen für ruhende Körper in eindeutiger Weise zu den Grundgleichungen für bewegte Körper.

Nämlich nach dem zweiten Axiom ist in jedem Raum-Zeitpunkte der Betrag des Geschwindigkeitsvektors $|\mathfrak{w}| < 1$. Infolgedessen können wir dem Vektor $\mathfrak{w}$ stets umkehrbar eindeutig das Quadrupel von Größen

$$w_1 = \frac{w_x}{\sqrt{1-w^2}}, \quad w_2 = \frac{w_y}{\sqrt{1-w^2}}, \quad w_3 = \frac{w_z}{\sqrt{1-w^2}}, \quad w_4 = \frac{i}{\sqrt{1-w^2}}$$

zuordnen, zwischen denen die Beziehung

$$(27) \quad w_1^2 + w_2^2 + w_3^2 + w_4^2 = -1$$

statthat. Aus den Ausführungen am Schlusse des § 4 ist ersichtlich, daß dieses Quadrupel sich bei Lorentz-Transformationen als Raum-Zeit-Vektor I. Art verhält, und wir wollen es den *Raum-Zeit-Vektor Geschwindigkeit* nennen.

Fassen wir nun eine bestimmte Stelle x, y, z der Materie zu einer bestimmten Zeit t auf. Ist in diesem Raum-Zeitpunkte $\mathfrak{w} = 0$, so haben wir für ihn nach dem ersten Axiom unmittelbar die Gleichungen (A), (B), (V) aus § 7. Ist in ihm $\mathfrak{w} \neq 0$, so existiert, weil $|\mathfrak{w}| < 1$ ist, nach (16) eine spezielle Lorentz-Transformation, deren Vektor $\mathfrak{v}$ gleich diesem Vektor $\mathfrak{w}(x, y, z, t)$ ist, und wir gehen allgemein zu einem neuen Bezugssystem x', y', z', t' gemäß dieser bestimmten Transformation über. Für den betrachteten Raum-Zeitpunkt entstehen dabei, wie wir in § 4 sahen, die neuen Werte

$$(28) \quad w_1' = 0, \quad w_2' = 0, \quad w_3' = 0, \quad w_4' = i,$$

und also der neue Geschwindigkeitsvektor $w' = 0$, der *Raum-Zeitpunkt* wird, wie wir uns dort ausdrückten, auf *Ruhe transformiert*. Nun sollen nach dem dritten Axiom aus den Grundgleichungen für den Raum-Zeitpunkt x, y, z, t dabei die Grundgleichungen für das entsprechende System x', y', z', t' , geschrieben in den transformierten Größen $w', \varrho', \mathfrak{s}', e', m', \mathfrak{E}', \mathfrak{M}'$ und deren Differentialquotienten nach x', y', z', t' hervorgehen. Diese letzteren Gleichungen aber müssen, nach dem ersten Axiom, weil jetzt $w' = 0$ ist, genau sein:

1) diejenigen Differentialgleichungen (A'), (B'), die aus (A) und (B) einfach dadurch hervorgehen, daß alle Buchstaben dort mit einem oberen Strich versehen werden,

2) die Gleichungen

$$(V') \quad e' = \varepsilon \mathfrak{E}', \quad \mathfrak{M}' = \mu m', \quad \mathfrak{s}' = \sigma \mathfrak{E}',$$

wobei ε, μ, σ Dielektrizitätskonstante, magnetische Permeabilität, Leitfähigkeit für das System x', y', z', t' , d. i. also im betrachteten Raum-Zeitpunkte x, y, z, t der Materie sind.

Jetzt gehen wir durch die reziproke Lorentz-Transformation rückwärts zu den ursprünglichen Variablen x, y, z, t und den Größen $w, \varrho, \mathfrak{s}, e, m, \mathfrak{E}, \mathfrak{M}$ und die Gleichungen, die wir dann aus den eben genannten erhalten, werden die von uns gesuchten allgemeinen Grundgleichungen für bewegte Körper sein.

Nun ist aus den Ausführungen in § 4 und § 5 zu ersehen, daß sowohl das Gleichungssystem (A) für sich wie das Gleichungssystem (B) für sich kovariant bei den Lorentz-Transformationen ist; d. h. die Gleichungen, die wir von (A'), (B') rückwärts erlangen, müssen genau gleichlauten mit den Gleichungen (A), (B), wie wir sie für ruhende Körper annahmen. Wir haben also als erstes Ergebnis:

Von den Grundgleichungen der Elektrodynamik für bewegte Körper lauten die Differentialgleichungen, geschrieben in ϱ und den Vektoren $\mathfrak{s}, e, m, \mathfrak{E}, \mathfrak{M}$, genau wie für ruhende Körper. Die Geschwindigkeit der Materie tritt in diesen Gleichungen noch nicht auf. In vektorieller Schreibweise sind diese Gleichungen also wieder

$$(I) \quad \text{curl } m - \frac{\partial e}{\partial t} = \mathfrak{s},$$

$$(II) \quad \text{div } e = \varrho,$$

$$(III) \quad \text{curl } \mathfrak{E} + \frac{\partial \mathfrak{M}}{\partial t} = 0,$$

$$(IV) \quad \text{div } \mathfrak{M} = 0.$$

Die Geschwindigkeit der Materie wird ausschließlich auf die Zusatz-

bedingungen verwiesen, welche den Einfluß der Materie auf Grund ihrer speziellen Konstanten ε , μ , σ charakterisieren. Transformieren wir jetzt diese Zusatzbedingungen (V') zurück auf die ursprünglichen Koordinaten x , y , z und die ursprüngliche Zeit t .

Nach den Formeln (15) in § 4 ist für die Richtung des Vektors w die Komponente von e' dieselbe wie von $e + [wm]$, die von m' dieselbe wie von $m - [we]$, für jede dazu senkrechte Richtung $\bar{w}$ aber ist die Komponente von e' bzw. m' gleich der entsprechenden Komponente von $e + [wm]$ bzw. von $m - [we]$, jedesmal multipliziert noch mit $\frac{1}{\sqrt{1-w^2}}$.

Andererseits werden $\mathfrak{E}'$ und $\mathfrak{M}'$ hier zu $\mathfrak{E} + [w\mathfrak{M}]$ und $\mathfrak{M} - [w\mathfrak{E}]$ in den ganz analogen Beziehungen stehen wie e' und m' zu $e + [wm]$ und $m - [we]$. So führt die Relation $e' = \varepsilon \mathfrak{E}'$, indem man bei den Vektoren zuerst die Komponenten nach der Richtung w , dann diejenigen nach zwei zu w und aufeinander senkrechten Richtungen $\bar{w}$ behandelt und die in letzteren Fällen entstehenden Gleichungen mit $\sqrt{1-w^2}$ multipliziert, zu

$$(C) \quad e + [wm] = \varepsilon (\mathfrak{E} + [w\mathfrak{M}]).$$

Die Relation $\mathfrak{M}' = \mu m'$ wird analog auf

$$(D) \quad \mathfrak{M} - [w\mathfrak{E}] = \mu (m - [we])$$

hinauslaufen.

Weiter folgt nach den Transformationsgleichungen (12), (10), (11) in § 4, indem dort q , r_v , $r_{\bar{v}}$, t , r'_v , $r'_{\bar{v}}$, t' durch $|w|$, $\mathfrak{s}_w$, $\mathfrak{s}_{\bar{w}}$, ϱ , $\mathfrak{s}'_w$, $\mathfrak{s}'_{\bar{w}}$, ϱ' zu ersetzen sind,

$$\varrho' = \frac{-|w|\mathfrak{s}_w + \varrho}{\sqrt{1-w^2}}, \quad \mathfrak{s}'_w = \frac{\mathfrak{s}_w - |w|\varrho}{\sqrt{1-w^2}}, \quad \mathfrak{s}'_{\bar{w}} = \mathfrak{s}_{\bar{w}},$$

so daß aus $\mathfrak{s}' = \sigma \mathfrak{E}'$ nunmehr

$$(E) \quad \frac{\mathfrak{s}_w - |w|\varrho}{\sqrt{1-w^2}} = \sigma (\mathfrak{E} + [w\mathfrak{M}])_w, \\ \mathfrak{s}_{\bar{w}} = \frac{\sigma (\mathfrak{E} + [w\mathfrak{M}])_{\bar{w}}}{\sqrt{1-w^2}}$$

hervorgeht. Nach der Art, wie hier die Leitfähigkeit σ eingeht, wird es angemessen sein, den Vektor $\mathfrak{s} - \varrho w$ mit den Komponenten $\mathfrak{s}_w - \varrho|w|$ nach der Richtung w und $\mathfrak{s}_{\bar{w}}$ nach den auf w senkrechten Richtungen $\bar{w}$, der für $\sigma = 0$ verschwindet, als *Leitungsstrom* zu bezeichnen.

Wir bemerken, daß für $\varepsilon = 1$, $\mu = 1$ die Gleichungen $e' = \mathfrak{E}'$, $m' = \mathfrak{M}'$ durch die reziproke Lorentz-Transformation, die hier die spezielle mit $-w$ als Vektor wird, gemäß (15) sofort zu $e = \mathfrak{E}$, $m = \mathfrak{M}$ führen und daß für $\sigma = 0$ die Gleichung $\mathfrak{s}' = 0$ zu $\mathfrak{s} = \varrho w$ führt, so daß in der Tat als Grenzfall der hier erhaltenen Gleichungen für $\varepsilon = 1$, $\mu = 1$, $\sigma = 0$ sich die in § 2 betrachteten „Grundgleichungen für den Äther“ ergeben.

§ 9.

Die Grundgleichungen in der Theorie von Lorentz.

Sehen wir nun zu, inwieweit die Grundgleichungen, die Lorentz annimmt, dem Relativitätspostulate, das soll heißen dem in § 8 formulierten Relativitätsprinzipie entsprechen. In dem Artikel „Elektronentheorie“ (Enzykl. der math. Wiss., V 2, Art. 14) hat Lorentz für beliebige, auch magnetisierte Körper zunächst die Differentialgleichungen (s. dort S. 209 unter Berücksichtigung von Gl. XXX' daselbst und von Formel (14) auf S. 78 desselben Heftes):

$$(IIIa'') \quad \text{curl} (\mathfrak{H} - [w \mathfrak{E}]) = \mathfrak{J} + \frac{\partial \mathfrak{D}}{\partial t} + w \text{div} \mathfrak{D} - \text{curl} [w \mathfrak{D}],$$

$$(I'') \quad \text{div} \mathfrak{D} = \varrho,$$

$$(IV'') \quad \text{curl} \mathfrak{E} = - \frac{\partial \mathfrak{B}}{\partial t},$$

$$(V'') \quad \text{div} \mathfrak{B} = 0.$$

Dann setzt Lorentz für bewegte nicht magnetisierte Körper (S. 223, Z. 3) $\mu = 1$, $\mathfrak{B} = \mathfrak{H}$ und nimmt dazu das Eingehen der Dielektrizitätskonstante ε und der Leitfähigkeit σ gemäß

$$(Gl. XXXIV'', S. 227) \quad \mathfrak{D} - \mathfrak{E} = (\varepsilon - 1) (\mathfrak{E} + [w \mathfrak{B}]),$$

$$(Gl. XXXIII'', S. 223) \quad \mathfrak{J} = \sigma (\mathfrak{E} + [w \mathfrak{B}])$$

an. Die Lorentzschen Zeichen $\mathfrak{E}$, $\mathfrak{B}$, $\mathfrak{D}$, $\mathfrak{H}$ sind hier durch $\mathfrak{E}$, $\mathfrak{M}$, $\mathfrak{e}$, $\mathfrak{m}$ ersetzt, während $\mathfrak{J}$ bei Lorentz als Leitungsstrom bezeichnet wird.

Die drei letzten der zitierten Differentialgleichungen nun decken sich sofort mit den Gleichungen (II), (III), (IV) hier, die erste Gleichung aber würde, indem wir $\mathfrak{J}$ mit dem für $\sigma = 0$ verschwindenden Strome $\mathfrak{s} - w\varrho$ identifizieren, in

$$(29) \quad \text{curl} (\mathfrak{H} - [w \mathfrak{E}]) = \mathfrak{s} + \frac{\partial \mathfrak{D}}{\partial t} - \text{curl} [w \mathfrak{D}]$$

übergehen und verschieden von (I) hier ausfallen. Danach entsprechen die allgemeinen Differentialgleichungen von Lorentz für beliebig magnetisierte Körper *nicht* dem Relativitätsprinzipie.

Andererseits würde die dem Relativitätsprinzipie entsprechende Form für die Bedingung des Nichtmagnetisiertseins aus (D) in § 8 mit $\mu = 1$ *nicht* wie bei Lorentz als $\mathfrak{B} = \mathfrak{H}$, sondern als

$$(30) \quad \mathfrak{B} - [w \mathfrak{E}] = \mathfrak{H} - [w \mathfrak{D}] \quad (\text{hier } \mathfrak{M} - [w \mathfrak{E}] = \mathfrak{m} - [w \mathfrak{e}])$$

anzunehmen sein. Nun geht aber die zuletzt hingeschriebene Differentialgleichung (29) durch $\mathfrak{H} = \mathfrak{B}$ in dieselbe Gleichung (abgesehen von der Verschiedenheit der Zeichen) über, in welche (I) hier sich durch

$m - [w\epsilon] = \mathfrak{M} - [w\mathfrak{E}]$ verwandeln würde. So kommt es durch eine Kompensation zweier Widersprüche gegen das Relativitätsprinzip zustande, daß für nicht magnetisierte bewegte Körper die Differentialgleichungen von Lorentz sich zuletzt dem Relativitätsprinzip doch anpassen.

Macht man weiter für nicht magnetisierte Körper von (30) hier Gebrauch und setzt demgemäß $\mathfrak{H} = \mathfrak{B} + [w, \mathfrak{D} - \mathfrak{E}]$, so würde zufolge (C) in § 8

$$(\epsilon - 1)(\mathfrak{E} + [w\mathfrak{B}]) = \mathfrak{D} - \mathfrak{E} + [w[w, \mathfrak{D} - \mathfrak{E}]]$$

anzunehmen sein, d. i. für die Richtung von w :

$$(\epsilon - 1)(\mathfrak{E} + [w\mathfrak{B}])_w = (\mathfrak{D} - \mathfrak{E})_w,$$

und für jede zu w senkrechte Richtung $\bar{w}$:

$$(\epsilon - 1)(\mathfrak{E} + [w\mathfrak{B}])_{\bar{w}} = (1 - w^2)(\mathfrak{D} - \mathfrak{E})_{\bar{w}},$$

d. i. mit der oben genannten Lorentzschen Annahme nur in Übereinstimmung bis auf Fehler von der Ordnung w^2 gegen 1.

Auch nur mit dem gleichen Grade der Annäherung entspricht der oben genannte Lorentzsche Ansatz für $\mathfrak{J}$ den durch das Relativitätsprinzip geforderten Beziehungen (vgl. (E) in § 8), daß die Komponenten $\mathfrak{J}_w$ bzw. $\mathfrak{J}_{\bar{w}}$ gleich den entsprechenden Komponenten von $\sigma(\mathfrak{E} + [w\mathfrak{B}])$, multipliziert in $\sqrt{1 - w^2}$ bzw. in $\frac{1}{\sqrt{1 - w^2}}$ seien.

§ 10.

Die Grundgleichungen nach E. Cohn.

E. Cohn*) nimmt folgende Grundgleichungen an:

$$(31) \quad \text{curl}(M + [w\mathfrak{E}]) = \frac{\partial \mathfrak{E}}{\partial t} + w \text{div } \mathfrak{E} + \mathfrak{J},$$

$$- \text{curl}(E - [w\mathfrak{M}]) = \frac{\partial \mathfrak{M}}{\partial t} + w \text{div } \mathfrak{M},$$

$$(32) \quad \mathfrak{J} = \sigma E, \quad \mathfrak{E} = \epsilon E - [wM], \quad \mathfrak{M} = \mu M + [wE],$$

wobei E, M als elektrische und magnetische Feldintensität (Kraft), $\mathfrak{E}, \mathfrak{M}$ als elektrische und magnetische Polarisierung (Erregung) aufgefaßt werden. Die Gleichungen lassen noch das Vorhandensein von wahren Magnetismus zu; wollen wir davon absehen, so ist $\text{div } \mathfrak{M} = 0$ zu setzen.

Ein Einwand gegen diese Gleichungen ist, daß nach ihnen für $\epsilon = 1$ $\mu = 1$ nicht die Vektoren Kraft und Erregung zusammenfallen. Fassen wir jedoch in den Gleichungen nicht E und M , sondern $E - [w\mathfrak{M}]$ und $M + [w\mathfrak{E}]$ als elektrische und magnetische Kraft auf und substituieren

*) Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen, mathematisch-physikalische Klasse, 1901, S. 74 (auch in Annalen der Physik, Bd. 7 (4), 1902, S. 29).

im Hinblick hierauf für $\mathfrak{E}$, $\mathfrak{M}$, E , M , $\text{div } \mathfrak{E}$ die Zeichen ϵ , $\mathfrak{M}$, $\mathfrak{E} + [w\mathfrak{M}]$, $m - [we]$, ϵ , so gehen zunächst die Differentialgleichungen in unsere Gleichungen über und zugleich verwandeln die Bedingungen (32) sich in

$$\begin{aligned}\mathfrak{S} &= \sigma(\mathfrak{E} + [w\mathfrak{M}]), \\ \epsilon + [w, m - [we]] &= \epsilon(\mathfrak{E} + [w\mathfrak{M}]), \\ \mathfrak{M} - [w, \mathfrak{E} + [w\mathfrak{M}]] &= \mu(m - [we]);\end{aligned}$$

damit würden in der Tat diese Gleichungen von Cohn bis auf Fehler von der Ordnung w^2 gegen 1 genau die durch das Relativitätsprinzip geforderten werden.

Erwähnt sei noch, daß die von Hertz angenommenen Gleichungen (in den Bezeichnungen von Cohn) lauten wie (31) mit den anderen Zusatzbedingungen

$$(33) \quad \mathfrak{E} = \epsilon E, \quad \mathfrak{M} = \mu M, \quad \mathfrak{S} = \sigma E;$$

und dieses Gleichungssystem würde auch nicht bei irgendwelcher veränderten Bezugnahme der Zeichen auf beobachtbare Größen sich dem Relativitätsprinzip bis auf Fehler von der Ordnung w^2 gegen 1 anpassen.

§ 11.

Typische Darstellung der Grundgleichungen.

Bei der Aufstellung der Grundgleichungen leitete uns der Gedanke, für sie eine Kovarianz bezüglich der Gruppe der Lorentz-Transformationen zu erzielen. Jetzt haben wir noch die ponderomotorischen Wirkungen und die Umsetzung der Energie im elektromagnetischen Felde zu behandeln, und da kann es von vornherein nicht zweifelhaft sein, daß die Erledigung dieser Fragen jedenfalls zusammenhängen wird mit den einfachsten, an die Grundgleichungen anknüpfenden Bildungen, die wieder Kovarianz bei den Lorentz-Transformationen zeigen. Um auf diese Bildungen hingewiesen zu werden, will ich vor allem die Grundgleichungen jetzt in eine *typische Form* bringen, die ihre Kovarianz bei der Lorentz-schen Gruppe in Evidenz setzt. Dabei bediene ich mich einer Rechnungsmethode, die ein abgekürztes Operieren mit den Raum-Zeit-Vektoren I. und II. Art bezweckt, und deren Regeln und Bezeichnungen, soweit sie für uns nützlich sein werden, ich hier zuvörderst zusammenstelle.

1°. Ein System von Größen

$$\begin{vmatrix} a_{11}, & \dots, & a_{1q} \\ \vdots & & \vdots \\ a_{p1}, & \dots, & a_{pq} \end{vmatrix},$$

angeordnet in p -Horizontal-, q -Vertikalreihen heißt eine $p \times q$ -reihige *Matrix**) und werde mit einem einzigen Zeichen, etwa hier A , bezeichnet.

Werden alle Größen a_{hk} mit dem nämlichen Faktor c multipliziert, so soll die entstehende Matrix der Größen ca_{hk} mit cA bezeichnet werden.

Werden die Rollen der Horizontal- und Vertikalreihen in A vertauscht, so erhält man eine $q \times p$ -reihige Matrix, welche die *transponierte* von A heißt und mit $\bar{A}$ bezeichnet werden soll:

$$\bar{A} = \begin{vmatrix} a_{11}, & \dots, & a_{p1} \\ \vdots & & \vdots \\ a_{1q}, & \dots, & a_{pq} \end{vmatrix}.$$

Hat man eine zweite Matrix mit gleichen Anzahlen p und q , wie A ,

$$B = \begin{vmatrix} b_{11}, & \dots, & b_{1q} \\ \vdots & & \vdots \\ a_{p1}, & \dots, & b_{pq} \end{vmatrix},$$

so soll $A + B$ die ebenfalls $p \times q$ -reihige Matrix aus den entsprechenden Binomen $a_{hk} + b_{hk}$ bedeuten.

2°. Hat man zwei Matrizen

$$A = \begin{vmatrix} a_{11}, & \dots, & a_{1q} \\ \vdots & & \vdots \\ a_{p1}, & \dots, & a_{pq} \end{vmatrix}, \quad B = \begin{vmatrix} b_{11}, & \dots, & b_{1r} \\ \vdots & & \vdots \\ b_{q1}, & \dots, & b_{qr} \end{vmatrix},$$

wobei die Anzahl der Horizontalreihen der zweiten gleich der Anzahl der Vertikalreihen der ersten ist, so wird unter AB , dem Produkte aus A und B , die Matrix

$$C = \begin{vmatrix} c_{11}, & \dots, & c_{1r} \\ \vdots & & \vdots \\ c_{p1}, & \dots, & c_{pr} \end{vmatrix}$$

verstanden, deren Elemente durch Kombination der Horizontalreihen von A und der Vertikalreihen von B nach der Regel

$$c_{hk} = a_{h1}b_{1k} + a_{h2}b_{2k} + \dots + a_{hq}b_{qk} \quad \begin{matrix} (h = 1, 2, \dots, p) \\ (k = 1, 2, \dots, r) \end{matrix}$$

gebildet sind. Für solche Produkte gilt das *assoziative* Gesetz $(AB)S = A(BS)$; hierbei ist unter S eine dritte Matrix gedacht mit so viel Horizontalreihen, als B (und damit auch AB) Vertikalreihen hat.

Für die transponierte Matrix zu $C = AB$ gilt $\bar{C} = \bar{B}\bar{A}$.

3°. Es werden hier nur Matrizen in Betracht kommen mit höchstens vier Horizontalreihen und höchstens vier Vertikalreihen.

*) Man könnte auch daran denken, statt des Cayleyschen Matrizenkalküls den Hamiltonschen Quaternionenkalkül heranzuziehen, doch erscheint mir der letztere für unsere Zwecke als zu eng und schwerfällig.

Als *Einheitsmatrix* (und in Gleichungen für Matrizen kurzweg mit 1) werde die 4×4 -reihige Matrix der folgenden Elemente

$$(34) \quad \begin{vmatrix} e_{11}, e_{12}, e_{13}, e_{14} \\ e_{21}, e_{22}, e_{23}, e_{24} \\ e_{31}, e_{32}, e_{33}, e_{34} \\ e_{41}, e_{42}, e_{43}, e_{44} \end{vmatrix} = \begin{vmatrix} 1, 0, 0, 0 \\ 0, 1, 0, 0 \\ 0, 0, 1, 0 \\ 0, 0, 0, 1 \end{vmatrix}$$

bezeichnet. Für ein Vielfaches $c \cdot 1$ der Einheitsmatrix (in dem unter 1^o festgesetzten Sinne einer Matrix cA) soll dann in Gleichungen für Matrizen kurzweg c stehen.

Für eine 4×4 -reihige Matrix A soll $\text{Det } A$ die Determinante aus den 4×4 Elementen der Matrix bedeuten. Ist dann $\text{Det } A \neq 0$, so gehört zu A eine bestimmte *reziproke* Matrix, mit A^{-1} bezeichnet, so daß $A^{-1} A = 1$ wird. —

Eine Matrix

$$f = \begin{vmatrix} 0, f_{12}, f_{13}, f_{14} \\ f_{21}, 0, f_{23}, f_{24} \\ f_{31}, f_{32}, 0, f_{34} \\ f_{41}, f_{42}, f_{43}, 0 \end{vmatrix},$$

in welcher die Elemente die Relationen $f_{kh} = -f_{hk}$ erfüllen, heißt eine *alternierende* Matrix. Diese Relationen besagen, daß die transponierte Matrix $\bar{f} = -f$ ist. Alsdann werde mit f^* und als die *duale* Matrix von f die ebenfalls alternierende Matrix

$$(35) \quad f^* = \begin{vmatrix} 0, f_{34}, f_{42}, f_{23} \\ f_{43}, 0, f_{14}, f_{31} \\ f_{24}, f_{41}, 0, f_{12} \\ f_{32}, f_{13}, f_{21}, 0 \end{vmatrix}$$

bezeichnet. Dabei wird

$$(36) \quad f^* f = f_{32} f_{14} + f_{13} f_{24} + f_{21} f_{34},$$

das soll nun heißen eine 4×4 -reihige Matrix, in der alle Elemente außerhalb der Hauptdiagonale von links oben nach rechts unten Null sind und alle Elemente in dieser Diagonale untereinander übereinstimmen und gleich der hier rechts genannten Verbindung aus den Koeffizienten von f sind. Die Determinante von f erweist sich dann als das Quadrat

dieser Verbindung und wir wollen das Zeichen $\text{Det } \frac{1}{2} f$ eindeutig als die *Abkürzung*

$$(37) \quad \text{Det } \frac{1}{2} f = f_{32} f_{14} + f_{13} f_{24} + f_{21} f_{34}$$

erklären.

4°. Eine lineare Transformation

$$(38) \quad x_h = \alpha_{h1} x_1' + \alpha_{h2} x_2' + \alpha_{h3} x_3' + \alpha_{h4} x_4' \quad (h = 1, 2, 3, 4)$$

werde auch einfach durch die 4×4 -reihige Matrix der Koeffizienten

$$A = \begin{vmatrix} \alpha_{11} & \alpha_{12} & \alpha_{13} & \alpha_{14} \\ \alpha_{21} & \alpha_{22} & \alpha_{23} & \alpha_{24} \\ \alpha_{31} & \alpha_{32} & \alpha_{33} & \alpha_{34} \\ \alpha_{41} & \alpha_{42} & \alpha_{43} & \alpha_{44} \end{vmatrix},$$

als Transformation A, bezeichnet. Durch die Transformation A geht der Ausdruck

$$x_1^2 + x_2^2 + x_3^2 + x_4^2$$

in die quadratische Form

$$\sum a_{hk} x_h' x_k' \quad (h, k = 1, 2, 3, 4)$$

über, wobei

$$a_{hk} = \alpha_{1h} \alpha_{1k} + \alpha_{2h} \alpha_{2k} + \alpha_{3h} \alpha_{3k} + \alpha_{4h} \alpha_{4k}$$

wird, d. h. die 4×4 -reihige (symmetrische) Matrix der Koeffizienten a_{hk} dieser Form wird das Produkt $\bar{A}A$ der transponierten Matrix von A in die Matrix A. Soll also durch die Transformation der neue Ausdruck

$$x_1'^2 + x_2'^2 + x_3'^2 + x_4'^2$$

hervorgehen, so muß

$$(39) \quad \bar{A}A = 1$$

die Matrix 1 werden. Dieser Relation hat demnach A zu entsprechen, wenn die Transformation (38) eine Lorentz-Transformation sein soll. Für die Determinante von A folgt aus (39): $(\text{Det } A)^2 = 1$, $\text{Det } A = \pm 1$. Die Bedingung (39) kommt zugleich auf

$$(40) \quad A^{-1} = \bar{A}$$

hinaus, d. h. die reziproke Matrix von A muß sich mit der transponierten von A decken.

Für A als Lorentz-Transformation haben wir noch weiter die Bestimmungen getroffen, daß $\text{Det } A = +1$ sei, daß jede der Größen α_{14} , α_{24} , α_{34} , α_{41} , α_{42} , α_{43} rein imaginär (bzw. Null), die anderen Koeffizienten in A reell seien und endlich noch $\alpha_{44} > 0$ sei.

5°. Ein Raum-Zeit-Vektor I. Art s_1, s_2, s_3, s_4 soll durch die 1×4 -reihige Matrix seiner vier Komponenten:

$$(41) \quad s = | s_1, s_2, s_3, s_4 |$$

repräsentiert werden und ist bei einer Lorentz-Transformation A durch sA zu ersetzen.

Ein Raum-Zeit-Vektor II. Art mit den *Komponenten* $f_{23}, f_{31}, f_{13}, f_{14}, f_{24}, f_{34}$ soll durch die alternierende Matrix

$$(42) \quad f = \begin{vmatrix} 0, & f_{12}, & f_{13}, & f_{14} \\ f_{21}, & 0, & f_{23}, & f_{24} \\ f_{31}, & f_{32}, & 0, & f_{34} \\ f_{41}, & f_{42}, & f_{43}, & 0 \end{vmatrix}$$

repräsentiert werden und ist (s. die in § 5 (23) und (24) festgesetzte Regel) bei einer Lorentz-Transformation A durch $\bar{A}fA = A^{-1}fA$ zu ersetzen. Dabei gilt in bezug auf den Ausdruck (37) die Identität $\text{Det}^{\frac{1}{2}}(\bar{A}fA) = \text{Det} A \text{Det}^{\frac{1}{2}}f$. Es wird danach $\text{Det}^{\frac{1}{2}}f$ eine Invariante bei den Lorentz-Transformationen (s. Gleichung (26) in § 5).

Für die duale Matrix f^* folgt dann mit Rücksicht auf (36):

$$(A^{-1}f^*A)(A^{-1}fA) = A^{-1}f^*fA = \text{Det}^{\frac{1}{2}}f \cdot A^{-1}A = \text{Det}^{\frac{1}{2}}f,$$

woraus zu ersehen ist, daß mit dem Raum-Zeit-Vektor II. Art f zusammen auch die zugehörige duale Matrix f^* sich wie ein Raum-Zeit-Vektor II. Art abändert, und es heiße deshalb f^* mit den Komponenten $f_{14}, f_{24}, f_{34}, f_{23}, f_{31}, f_{12}$ der *duale Raum-Zeit-Vektor* von f .

6°. Sind w und s zwei Raum-Zeit-Vektoren I. Art, so wird unter $w\bar{s}$ (wie auch unter $s\bar{w}$) die Verbindung

$$(43) \quad w_1s_1 + w_2s_2 + w_3s_3 + w_4s_4$$

aus den bezüglichen Komponenten zu verstehen sein. Bei einer Lorentz-Transformation A ist wegen $(wA)(\bar{A}s) = w\bar{s}$ diese Verbindung invariant. — Ist $w\bar{s} = 0$, so sollen w und s *normal* zueinander heißen.

Zwei Raum-Zeit-Vektoren I. Art w, s geben ferner zur Bildung der 2×4 -reihigen Matrix

$$\begin{vmatrix} w_1, & w_2, & w_3, & w_4 \\ s_1, & s_2, & s_3, & s_4 \end{vmatrix}$$

Anlaß. Es zeigt sich dann sofort, daß das System der sechs Größen

$$(44) \quad \begin{aligned} &w_2s_3 - w_3s_2, \quad w_3s_1 - w_1s_3, \quad w_1s_2 - w_2s_1, \quad w_1s_4 - w_4s_1, \\ &w_2s_4 - w_4s_2, \quad w_3s_4 - w_4s_3 \end{aligned}$$

sich bei den Lorentz-Transformationen als Raum-Zeit-Vektor II. Art verhält. Der Vektor II. Art mit diesen Komponenten (44) werde mit $[w, s]$ bezeichnet. Man erschließt leicht $\text{Det}^{\frac{1}{2}}[w, s] = 0$. Der duale Vektor von $[w, s]$ soll $[w, s]^*$ geschrieben werden.

Ist w ein Raum-Zeit-Vektor I. Art, f ein Raum-Zeit-Vektor II. Art, so bedeutet wf zunächst jedenfalls eine 1×4 -reihige Matrix. Bei einer Lorentz-Transformation A geht w in $w' = wA$, f in $f' = A^{-1}fA$ über; dabei wird $w'f' = wAA^{-1}fA = (wf)A$, d. h. wf transformiert sich wieder als ein Raum-Zeit-Vektor I. Art.

Man verifiziert, wenn w ein Vektor I., f ein Vektor II. Art ist, leicht die wichtige Identität

$$(45) \quad [w, wf] + [w, wf^*]^* = (w\bar{w})f.$$

Die Summe der zwei Raum-Zeit-Vektoren II. Art links ist im Sinne der Summe zweier alternierenden Matrizen zu verstehen.

Nämlich für $w_1 = 0$, $w_2 = 0$, $w_3 = 0$, $w_4 = i$ wird

$$wf = |if_{41}, if_{42}, if_{43}, 0|; \quad wf^* = |if_{32}, if_{13}, if_{21}, 0|;$$

$$[w, wf] = 0, 0, 0, f_{41}, f_{42}, f_{43}; \quad [w, wf^*] = 0, 0, 0, f_{32}, f_{13}, f_{21},$$

und die Bemerkung, daß in diesem speziellen Falle die Relation (45) zutrifft, genügt bereits, um derselben allgemein sicher zu sein, da diese Relation kovarianten Charakter für die Lorentz-Gruppe hat und zudem in w_1, w_2, w_3, w_4 homogen ist.

Nach diesen Vorbereitungen beschäftigen wir uns zunächst mit den Gleichungen (C), (D), (E), durch welche die Konstanten ε, μ, σ eingeführt werden.

Statt des Raumvektors w , Geschwindigkeit der Materie, führen wir, wie schon in § 8, den Raum-Zeit-Vektor I. Art w mit den vier Komponenten

$$w_1 = \frac{w_x}{\sqrt{1-w^2}}, \quad w_2 = \frac{w_y}{\sqrt{1-w^2}}, \quad w_3 = \frac{w_z}{\sqrt{1-w^2}}, \quad w_4 = \frac{i}{\sqrt{1-w^2}}$$

ein; dabei gilt

$$(46) \quad w\bar{w} = w_1^2 + w_2^2 + w_3^2 + w_4^2 = -1$$

und $-iw_4 > 0$.

Unter F und f wollen wir jetzt wieder die in den Grundgleichungen auftretenden Raum-Zeit-Vektoren II. Art $\mathfrak{M}$, $-i\mathfrak{E}$ und m , $-ie$ verstehen.

In $\Phi = -wF$ haben wir wieder einen Raum-Zeit-Vektor I. Art; seine Komponenten werden sein

$$\begin{aligned} \Phi_1 &= w_2 F_{12} + w_3 F_{13} + w_4 F_{14}, \\ \Phi_2 &= w_1 F_{21} + w_3 F_{23} + w_4 F_{24}, \\ \Phi_3 &= w_1 F_{31} + w_2 F_{32} + w_4 F_{34}, \\ \Phi_4 &= w_1 F_{41} + w_2 F_{42} + w_3 F_{43} \end{aligned}$$

Die drei ersten Größen Φ_1, Φ_2, Φ_3 sind bzw. die x -, y -, z -Komponente des Raumvektors

$$(47) \quad \frac{\mathfrak{E} + [\mathfrak{w} \mathfrak{M}]}{\sqrt{1 - \mathfrak{w}^2}},$$

und ferner ist

$$(48) \quad \Phi_4 = \frac{i(\mathfrak{w} \mathfrak{E})}{\sqrt{1 - \mathfrak{w}^2}}.$$

Da die Matrix F eine alternierende ist, gilt offenbar

$$(49) \quad w\bar{\Phi} = w_1\Phi_1 + w_2\Phi_2 + w_3\Phi_3 + w_4\Phi_4 = 0,$$

der Vektor Φ ist also normal zu w ; wir können diese Relation auch schreiben:

$$(50) \quad \Phi_4 = i(\mathfrak{w}_x\Phi_1 + \mathfrak{w}_y\Phi_2 + \mathfrak{w}_z\Phi_3).$$

Den Raum-Zeit-Vektor I. Art Φ will ich *elektrische Ruh-Kraft* nennen.

Analoge Beziehungen wie zwischen $-wF, \mathfrak{E}, \mathfrak{M}, \mathfrak{w}$ stellen sich zwischen $-wf, \mathfrak{e}, \mathfrak{m}, \mathfrak{w}$ heraus und insbesondere wird auch $-wf$ normal zu w sein. Es kann nunmehr die Relation (C) durch

$$\{C\} \quad wf = \varepsilon wF$$

ersetzt werden, eine Formel, die zwar vier Gleichungen für die bezüglichen Komponenten liefert, jedoch so, daß die vierte im Hinblick auf (50) eine Folge der drei ersten ist.

Wir bilden ferner den Raum-Zeit-Vektor I. Art $\Psi = iw f^*$, dessen Komponenten sind:

$$\begin{aligned} \Psi_1 &= -i(w_2f_{34} + w_3f_{42} + w_4f_{23}), \\ \Psi_2 &= -i(w_1f_{43} + w_3f_{14} + w_4f_{31}), \\ \Psi_3 &= -i(w_1f_{24} + w_2f_{41} + w_4f_{12}), \\ \Psi_4 &= -i(w_1f_{32} + w_2f_{13} + w_3f_{21}). \end{aligned}$$

Davon sind die drei ersteren Ψ_1, Ψ_2, Ψ_3 bzw. die x -, y -, z -Komponente des Raumvektors

$$(51) \quad \frac{\mathfrak{m} - [\mathfrak{w} \mathfrak{e}]}{\sqrt{1 - \mathfrak{w}^2}},$$

und weiter ist

$$(52) \quad \Psi_4 = \frac{i(\mathfrak{w} \mathfrak{m})}{\sqrt{1 - \mathfrak{w}^2}};$$

zwischen ihnen besteht die Beziehung

$$(53) \quad w\bar{\Psi} = w_1\Psi_1 + w_2\Psi_2 + w_3\Psi_3 + w_4\Psi_4 = 0,$$

die wir auch

$$(54) \quad \Psi_4 = i(\mathfrak{w}_x\Psi_1 + \mathfrak{w}_y\Psi_2 + \mathfrak{w}_z\Psi_3)$$

schreiben können; der Vektor Ψ ist also wieder normal zu w . Den Raum-Zeit-Vektor I. Art Ψ will ich *magnetische Ruh-Kraft* nennen.

Analoge Beziehungen wie zwischen iwf^* , m , e , w haben zwischen iwF^* , M , $\mathfrak{E}$, w statt und es kann die Relation (D) nunmehr durch

$$\{D\} \quad wF^* = \mu wf^*$$

ersetzt werden.

Die Gleichungen $\{C\}$ und $\{D\}$ können wir benutzen, um die Feldvektoren F und f auf Φ und Ψ zurückzuführen. Wir haben

$$wF = -\Phi, \quad wF^* = -i\mu\Psi, \quad wf = -\varepsilon\Phi, \quad wf^* = -i\Psi$$

und die Anwendung der Regel (45) führt im Hinblick auf (46) zu

$$(55) \quad F = [w, \Phi] + i\mu[w, \Psi]^*,$$

$$(56) \quad f = \varepsilon[w, \Phi] + i[w, \Psi]^*,$$

d. i.

$$F_{12} = (w_1\Phi_2 - w_2\Phi_1) + i\mu(w_3\Psi_4 - w_4\Psi_3), \text{ u. s. f.}$$

$$f_{12} = \varepsilon(w_1\Phi_2 - w_2\Phi_1) + i(w_3\Psi_4 - w_4\Psi_3), \text{ u. s. f.}$$

Wir ziehen ferner den Raum-Zeit-Vektor II. Art $[\Phi, \Psi]$ mit den sechs Komponenten

$$\begin{aligned} \Phi_2\Psi_3 - \Phi_3\Psi_2, \quad \Phi_3\Psi_1 - \Phi_1\Psi_3, \quad \Phi_1\Psi_2 - \Phi_2\Psi_1, \\ \Phi_1\Psi_4 - \Phi_4\Psi_1, \quad \Phi_2\Psi_4 - \Phi_4\Psi_2, \quad \Phi_3\Psi_4 - \Phi_4\Psi_3 \end{aligned}$$

in Betracht. Alsdann verschwindet der zugehörige Raum-Zeit-Vektor I. Art

$$w[\Phi, \Psi] = - (w\bar{\Psi})\Phi + (w\bar{\Phi})\Psi$$

wegen (49) und (53) identisch. Führen wir nun den Raum-Zeit-Vektor I. Art

$$(57) \quad \Omega = iw[\Phi, \Psi]^*$$

mit den Komponenten

$$\Omega_1 = -i \begin{vmatrix} w_2, & w_3, & w_4 \\ \Phi_2, & \Phi_3, & \Phi_4 \\ \Psi_2, & \Psi_3, & \Psi_4 \end{vmatrix}, \text{ u. s. f.}$$

ein, so folgt durch Anwendung der Regel (45):

$$(58) \quad [\Phi, \Psi] = i[w, \Omega]^*,$$

d. i.

$$\Phi_1\Psi_2 - \Phi_2\Psi_1 = i(w_3\Omega_4 - w_4\Omega_3), \text{ u. s. f.}$$

Der Vektor Ω erfüllt offenbar die Relation

$$(59) \quad (w\bar{\Omega}) = w_1\Omega_1 + w_2\Omega_2 + w_3\Omega_3 + w_4\Omega_4 = 0,$$

die wir auch

$$\Omega_4 = i(w_x\Omega_1 + w_y\Omega_2 + w_z\Omega_3)$$

schreiben können, ist also wieder *normal* zu w . Falls $w = 0$ ist, hat man $\Phi_4 = 0$, $\Psi_4 = 0$, $\Omega_4 = 0$ und

$$(60) \quad \Omega_1 = \Phi_2 \Psi_3 - \Phi_3 \Psi_2, \quad \Omega_2 = \Phi_3 \Psi_1 - \Phi_1 \Psi_3, \quad \Omega_3 = \Phi_1 \Psi_2 - \Phi_2 \Psi_1.$$

Den Raum-Zeit-Vektor I. Art Ω will ich als *Ruh-Strahl* bezeichnen.

Was die Relation (E) anbelangt, welche die Leitfähigkeit σ einführt, so erkennen wir zunächst, daß

$$-w\bar{s} = -(w_1 s_1 + w_2 s_2 + w_3 s_3 + w_4 s_4) = \frac{-|w| |\bar{s}_w + e|}{\sqrt{1-w^2}} = e'$$

die *Ruh-Dichte* der Elektrizität (s. § 8 und § 4 am Schlusse) wird. Als dann stellt

$$(61) \quad s + (w\bar{s})w$$

einen Raum-Zeit-Vektor I. Art vor, der wegen $w\bar{w} = -1$ offenbar wieder *normal* zu w ist und den ich als *Ruh-Strom* bezeichnen will. Fassen wir die drei ersten Komponenten dieses Vektors als x -, y -, z -Komponente eines Raum-Vektors auf, so ist für den letzteren die Komponente nach der Richtung von w :

$$\bar{s}_w - \frac{|w| |e'|}{\sqrt{1-w^2}} = \frac{\bar{s}_w - |w| |e|}{1-w^2} = \frac{\bar{S}_w}{1-w^2}$$

und die Komponente nach einer jeden zu w senkrechten Richtung $\bar{w}$ wieder

$$\bar{s}_{\bar{w}} = \bar{S}_{\bar{w}};$$

es hängt dieser Raum-Vektor also sehr einfach mit dem Raum-Vektor $\bar{S} = \bar{s} - e'w$ zusammen, den wir in § 8 als Leitungsstrom bezeichneten.

Nunmehr kann durch Vergleich mit $\Phi = -wF$ die Relation (E) auf die Gestalt gebracht werden:

$$\{E\} \quad s + (w\bar{s})w = -\sigma wF.$$

Diese Formel faßt wieder vier Gleichungen zusammen, von denen jedoch, weil es sich beiderseits um zu w normale Raum-Zeit-Vektoren I. Art handelt, die vierte eine Folge der drei ersten ist.

Endlich werden wir noch die Differentialgleichungen (A) und (B) in eine typische Form umsetzen.

§ 12.

Der Differentialoperator lor.

Eine 4×4 -reihige Matrix

$$(62) \quad S = \begin{vmatrix} S_{11} & S_{12} & S_{13} & S_{14} \\ S_{21} & S_{22} & S_{23} & S_{24} \\ S_{31} & S_{32} & S_{33} & S_{34} \\ S_{41} & S_{42} & S_{43} & S_{44} \end{vmatrix} = |S_{hk}|$$

mit der Vorschrift, sie bei einer Lorentz-Transformation A jedesmal durch $\bar{A}SA$ zu ersetzen, mag eine *Raum-Zeit-Matrix* II. Art heißen. Eine derartige Matrix hat man insbesondere

in der alternierenden Matrix f , die einem Raum-Zeit-Vektor II. Art f entspricht,

in dem Produkte fF zweier solcher alternierender Matrizen f, F , das bei einer Transformation A durch $(A^{-1}fA)(A^{-1}FA) = A^{-1}fFA$ zu ersetzen ist,

ferner, wenn w_1, w_2, w_3, w_4 und $\Omega_1, \Omega_2, \Omega_3, \Omega_4$ zwei Raum-Zeit-Vektoren I. Art sind, in der Matrix der 4×4 Elemente $S_{hk} = w_h \Omega_k$,

endlich in einem Vielfachen L der Einheitsmatrix, d. h. einer 4×4 -reihigen Matrix, in der alle Elemente in der Hauptdiagonale einen gleichen Wert L haben und die übrigen Elemente sämtlich Null sind.

Wir haben es hier stets mit Funktionen von Raum-Zeitpunkten x, y, z, it zu tun und können mit Vorteil eine 1×4 -reihige Matrix, gebildet aus den Differentiationssymbolen

$$\left| \frac{\partial}{\partial x}, \frac{\partial}{\partial y}, \frac{\partial}{\partial z}, \frac{\partial}{\partial it} \right|,$$

oder auch

$$(63) \quad \left| \frac{\partial}{\partial x_1}, \frac{\partial}{\partial x_2}, \frac{\partial}{\partial x_3}, \frac{\partial}{\partial x_4} \right|$$

geschrieben, verwenden. Für diese Matrix will ich die *Abkürzung* lor brauchen.

Es soll dann, wenn S wie in (62) eine Raum-Zeit-Matrix II. Art bedeutet, in sinngemäßer Übertragung der Regel für die Produktbildung von Matrizen, unter $\text{lor } S$ die 1×4 -reihige Matrix

$$|K_1, K_2, K_3, K_4|$$

der Ausdrücke

$$(64) \quad K_k = \frac{\partial S_{1k}}{\partial x_1} + \frac{\partial S_{2k}}{\partial x_2} + \frac{\partial S_{3k}}{\partial x_3} + \frac{\partial S_{4k}}{\partial x_4} \quad (k = 1, 2, 3, 4)$$

verstanden werden.

Wird durch eine Lorentz-Transformation A ein neues Bezugssystem x'_1, x'_2, x'_3, x'_4 für die Raum-Zeitpunkte eingeführt, so mag analog der Operator

$$\text{lor}' = \left| \frac{\partial}{\partial x'_1}, \frac{\partial}{\partial x'_2}, \frac{\partial}{\partial x'_3}, \frac{\partial}{\partial x'_4} \right|$$

angewandt werden. Geht dabei S in $S' = \bar{A}SA = |S'_{hk}|$ über, so wird dann unter $\text{lor}' S'$ die 1×4 -reihige Matrix der Ausdrücke

$$K'_k = \frac{\partial S'_{1k}}{\partial x'_1} + \frac{\partial S'_{2k}}{\partial x'_2} + \frac{\partial S'_{3k}}{\partial x'_3} + \frac{\partial S'_{4k}}{\partial x'_4} \quad (k = 1, 2, 3, 4)$$

zu verstehen sein. Nun gilt für die Differentiation einer beliebigen Funktion von einem Raum-Zeitpunkte die Regel

$$\begin{aligned}\frac{\partial}{\partial x'_k} &= \frac{\partial}{\partial x_1} \frac{\partial x_1}{\partial x'_k} + \frac{\partial}{\partial x_2} \frac{\partial x_2}{\partial x'_k} + \frac{\partial}{\partial x_3} \frac{\partial x_3}{\partial x'_k} + \frac{\partial}{\partial x_4} \frac{\partial x_4}{\partial x'_k} \\ &= \frac{\partial}{\partial x_1} \alpha_{1k} + \frac{\partial}{\partial x_2} \alpha_{2k} + \frac{\partial}{\partial x_3} \alpha_{3k} + \frac{\partial}{\partial x_4} \alpha_{4k},\end{aligned}$$

die in einer leicht verständlichen Weise symbolisch als

$$\text{lor}' = \text{lor } A$$

zu deuten ist, und mit Rücksicht hierauf folgt sogleich

$$(65) \quad \text{lor}' S' = \text{lor} (A(A^{-1}SA)) = (\text{lor } S)A,$$

d. h. wenn S eine Raum-Zeit-Matrix II. Art vorstellt, so transformiert sich $\text{lor } S$ als ein Raum-Zeit-Vektor I. Art.

Ist insbesondere L ein Vielfaches der Einheitsmatrix, so wird unter $\text{lor } L$ die Matrix der Elemente

$$(66) \quad \left| \frac{\partial L}{\partial x_1}, \frac{\partial L}{\partial x_2}, \frac{\partial L}{\partial x_3}, \frac{\partial L}{\partial x_4} \right|$$

zu verstehen sein.

Stellt $s = |s_1, s_2, s_3, s_4|$ einen Raum-Zeit-Vektor I. Art vor, so wird

$$(67) \quad \text{lor } \bar{s} = \frac{\partial s_1}{\partial x_1} + \frac{\partial s_2}{\partial x_2} + \frac{\partial s_3}{\partial x_3} + \frac{\partial s_4}{\partial x_4}$$

zu erklären sein. Treten bei Anwendung einer Lorentz-Transformation A die Zeichen lor' , s' an Stelle von lor , s , so folgt

$$\text{lor}' \bar{s}' = (\text{lor } A) (\bar{A}\bar{s}) = \text{lor } \bar{s},$$

d. h. $\text{lor } \bar{s}$ ist eine Invariante bei den Lorentz-Transformationen.

In allen diesen Beziehungen spielt der Operator lor selbst die Rolle eines Raum-Zeit-Vektors I. Art.

Stellt f einen Raum-Zeit-Vektor II. Art vor, so hat nun $-\text{lor } f$ den Raum-Zeit-Vektor I. Art mit den Komponenten

$$\begin{aligned}& \frac{\partial f_{12}}{\partial x_2} + \frac{\partial f_{13}}{\partial x_3} + \frac{\partial f_{14}}{\partial x_4}, \\ & \frac{\partial f_{21}}{\partial x_1} + \frac{\partial f_{23}}{\partial x_3} + \frac{\partial f_{24}}{\partial x_4}, \\ & \frac{\partial f_{31}}{\partial x_1} + \frac{\partial f_{32}}{\partial x_2} + \frac{\partial f_{34}}{\partial x_4}, \\ & \frac{\partial f_{41}}{\partial x_1} + \frac{\partial f_{42}}{\partial x_2} + \frac{\partial f_{43}}{\partial x_3}\end{aligned}$$

zu bedeuten. Hiernach läßt sich das System der Differentialgleichungen (A) in der kurzen Form

$$\{A\} \quad \text{lor } f = -s$$

zusammenziehen. Ganz entsprechend wird das System der Differentialgleichungen (B) zu schreiben sein:

$$\{B\} \quad \text{lor } F^* = 0.$$

Die im Hinblick auf die Definition (67) von $\text{lor } \bar{s}$ gebildeten Verbindungen $\text{lor } (\text{lor } f)$ und $\text{lor } (\text{lor } F^*)$ verschwinden offenbar identisch, indem f und F^* alternierende Matrizen sind. Darnach folgt aus $\{A\}$ für den Strom s die Beziehung

$$(68) \quad \frac{\partial s_1}{\partial x_1} + \frac{\partial s_2}{\partial x_2} + \frac{\partial s_3}{\partial x_3} + \frac{\partial s_4}{\partial x_4} = 0,$$

während die Relation

$$(69) \quad \text{lor } (\text{lor } \overline{F^*}) = 0$$

den Sinn hat, daß die vier in $\{B\}$ angewiesenen Gleichungen nur drei unabhängige Bedingungen für den Verlauf der Feldvektoren repräsentieren.

Ich fasse nunmehr die Resultate zusammen:

Es bedeute w den Raum-Zeit-Vektor I. Art $\frac{w}{\sqrt{1-w^2}}$, $\frac{i}{\sqrt{1-w^2}}$ (w Geschwindigkeit der Materie), F den Raum-Zeit-Vektor II. Art $\mathfrak{M}$, $-i\mathfrak{E}$ ($\mathfrak{M}$ magnetische Erregung, $\mathfrak{E}$ elektrische Kraft), f den Raum-Zeit-Vektor II. Art m , $-ie$ (m magnetische Kraft, e elektrische Erregung), s den Raum-Zeit-Vektor I. Art $\bar{s}$, $i\rho$ (ρ elektrische Raumdichte, $\bar{s} - \rho w$ Leitungsstrom), ε die Dielektrizitätskonstante, μ die magnetische Permeabilität, σ die Leitfähigkeit, so lauten (mit den in § 10 und § 11 erklärten Symbolen der Matrizenrechnung) die Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern

$$\begin{aligned} \{A\} \quad & \text{lor } f = -s, \\ \{B\} \quad & \text{lor } F^* = 0, \\ \{C\} \quad & wf = \varepsilon wF, \\ \{D\} \quad & wF^* = \mu wf^*, \\ \{E\} \quad & s + (w\bar{s})w = -\sigma wF. \end{aligned}$$

Dabei gilt $w\bar{w} = -1$, es sind die Raum-Zeit-Vektoren I. Art wF , wf , wF^* , wf^* , $s + (w\bar{s})w$ sämtlich normal zu w und endlich besteht für das Gleichungssystem $\{B\}$ der Zusammenhang

$$\text{lor } (\text{lor } \overline{F^*}) = 0.$$

In Anbetracht der zuletzt genannten Umstände steht hier genau die erforderliche Anzahl von unabhängigen Gleichungen zur Verfügung, um bei den geeigneten Grenzdaten die Vorgänge vollständig zu beschreiben, wofern die Bewegung der Materie, also der Vektor w als Funktion von x, y, z, t bekannt ist.

§ 13.

Das Produkt der Feldvektoren fF .

Endlich fragen wir nach den Gesetzen, die zur Bestimmung des Vektors w als Funktion von x, y, z, t führen. Bei den hierauf bezüglichen Untersuchungen treten diejenigen Ausdrücke in den Vordergrund, die durch *Bildung des Produkts der zwei alternierenden Matrizen*

$$f = \begin{vmatrix} 0, & f_{12}, & f_{13}, & f_{14} \\ f_{21}, & 0, & f_{23}, & f_{24} \\ f_{31}, & f_{32}, & 0, & f_{34} \\ f_{41}, & f_{42}, & f_{43}, & 0 \end{vmatrix}, \quad F = \begin{vmatrix} 0, & F_{12}, & F_{13}, & F_{14} \\ F_{21}, & 0, & F_{23}, & F_{24} \\ F_{31}, & F_{32}, & 0, & F_{34} \\ F_{41}, & F_{42}, & F_{43}, & 0 \end{vmatrix}$$

sich darbieten. Ich schreibe

$$(70) \quad fF = \begin{vmatrix} S_{11} - L, & S_{12}, & S_{13}, & S_{14} \\ S_{21}, & S_{22} - L, & S_{23}, & S_{24} \\ S_{31}, & S_{32}, & S_{33} - L, & S_{34} \\ S_{41}, & S_{42}, & S_{43}, & S_{44} - L \end{vmatrix}$$

so, daß dabei

$$(71) \quad S_{11} + S_{22} + S_{33} + S_{44} = 0$$

wird.

Alsdann bedeutet L die in den Indizes 1, 2, 3, 4 symmetrische Verbindung

$$(72) \quad L = \frac{1}{2} (f_{23}F_{23} + f_{31}F_{31} + f_{12}F_{12} + f_{14}F_{14} + f_{24}F_{24} + f_{34}F_{34}),$$

und es wird

$$(73) \quad \begin{aligned} S_{11} &= \frac{1}{2} (f_{23}F_{23} + f_{34}F_{34} + f_{42}F_{42} - f_{12}F_{12} - f_{13}F_{13} - f_{14}F_{14}), \\ S_{12} &= f_{13}F_{32} + f_{14}F_{42}, \text{ u. s. f.} \end{aligned}$$

Indem ich die *Realitätsverhältnisse* zum Ausdruck bringe, will ich noch

$$(74) \quad S = \begin{vmatrix} S_{11}, & S_{12}, & S_{13}, & S_{14} \\ S_{21}, & S_{22}, & S_{23}, & S_{24} \\ S_{31}, & S_{32}, & S_{33}, & S_{34} \\ S_{41}, & S_{42}, & S_{43}, & S_{44} \end{vmatrix} = \begin{vmatrix} X_x, & Y_x, & Z_x, & -iT_x \\ X_y, & Y_y, & Z_y, & -iT_y \\ X_z, & Y_z, & Z_z, & -iT_z \\ -iX_t, & -iY_t, & -iZ_t, & T_t \end{vmatrix}$$

schreiben, wobei dann

$$\begin{aligned}
 X_x &= \frac{1}{2} (m_x \mathfrak{M}_x - m_y \mathfrak{M}_y - m_z \mathfrak{M}_z + e_x \mathfrak{E}_x - e_y \mathfrak{E}_y - e_z \mathfrak{E}_z) \\
 X_y &= m_x \mathfrak{M}_y + e_y \mathfrak{E}_x, \quad Y_x = m_y \mathfrak{M}_x + e_x \mathfrak{E}_y, \quad \text{u. s. f.} \\
 X_z &= e_y \mathfrak{M}_z - e_z \mathfrak{M}_y, \\
 T_x &= m_z \mathfrak{E}_y - m_y \mathfrak{E}_z, \quad \text{u. s. f.} \\
 T_t &= \frac{1}{2} (m_x \mathfrak{M}_x + m_y \mathfrak{M}_y + m_z \mathfrak{M}_z + e_x \mathfrak{E}_x + e_y \mathfrak{E}_y + e_z \mathfrak{E}_z)
 \end{aligned}
 \tag{75}$$

und auch

$$L = \frac{1}{2} (m_x \mathfrak{M}_x + m_y \mathfrak{M}_y + m_z \mathfrak{M}_z - e_x \mathfrak{E}_x - e_y \mathfrak{E}_y - e_z \mathfrak{E}_z)
 \tag{76}$$

sämtlich reell sind. In den Theorien für ruhende Körper kommen die Verbindungen $X_x, X_y, X_z, Y_x, Y_y, Y_z, Z_x, Z_y, Z_z$ unter dem Namen „Maxwellsche Spannungen“, die Größen T_x, T_y, T_z als „Poyntingscher Vektor“, T_t als „elektromagnetische Energiedichte für die Volumeneinheit“ vor und wird L als „Lagrangesche Funktion“ bezeichnet.

Wir finden nun andererseits durch Zusammensetzung der zu f und F dualen Matrizen in umgekehrter Folge sofort

$$F^* f^* = \begin{vmatrix} -S_{11} - L, & -S_{12}, & -S_{13}, & -S_{14} \\ -S_{21}, & -S_{22} - L, & -S_{23}, & -S_{24} \\ -S_{31}, & -S_{32}, & -S_{33} - L, & -S_{34} \\ -S_{41}, & -S_{42}, & -S_{43}, & -S_{44} - L \end{vmatrix}
 \tag{77}$$

und können hiernach setzen

$$fF = S - L, \quad F^* f^* = -S - L,
 \tag{78}$$

indem wir unter L das Vielfache $L \cdot 1$ der Einheitsmatrix, d. h. die Matrix der Elemente

$$|L e_{hk}| \quad \left(\begin{matrix} e_{hh} = 1, & e_{hk} = 0, & h \neq k \\ h, k = 1, 2, 3, 4 \end{matrix} \right)$$

verstehen.

Daraus folgern wir weiter, indem hier $SL = LS$ ist,

$$F^* f^* f F = (-S - L)(S - L) = -SS + L^2,$$

und finden, da $f^* f = \text{Det}^{\frac{1}{2}} f$, $F^* F = \text{Det}^{\frac{1}{2}} F$ ist, die interessante Beziehung:

$$SS = L^2 - \text{Det}^{\frac{1}{2}} f \text{Det}^{\frac{1}{2}} F,
 \tag{79}$$

d. h. das Produkt der Matrix S in sich selbst ist ein Vielfaches der Einheitsmatrix, eine Matrix, in welcher außerhalb der Hauptdiagonale alle Elemente Null und in der Diagonale alle Elemente gleich sind und als gemeinsamen Wert die hier rechts angegebene Größe haben. Es gelten also allgemein die Relationen

$$(80) \quad S_{h1}S_{1k} + S_{h2}S_{2k} + S_{h3}S_{3k} + S_{h4}S_{4k} = 0$$

bei ungleichen Indizes h, k aus der Reihe 1, 2, 3, 4 und

$$(81) \quad S_{h1}S_{1h} + S_{h2}S_{2h} + S_{h3}S_{3h} + S_{h4}S_{4h} = L^2 - \text{Det}^{\frac{1}{2}} f \text{Det}^{\frac{1}{2}} F$$

für $h = 1, 2, 3, 4$.

Indem wir jetzt anstatt F und f in den Verbindungen (72), (73) mittels (55), (56), (57) die elektrische Ruh-Kraft Φ , die magnetische Ruh-Kraft Ψ , den Ruh-Strahl Ω einführen, gelangen wir zu den Ausdrücken:

$$(82) \quad L = -\frac{1}{2} \varepsilon \Phi \bar{\Phi} + \frac{1}{2} \mu \Psi \bar{\Psi},$$

$$(83) \quad S_{hk} = -\frac{1}{2} \varepsilon \Phi \bar{\Phi} e_{hk} - \frac{1}{2} \mu \Psi \bar{\Psi} e_{hk} \\ + \varepsilon (\Phi_h \Phi_k - \Phi \bar{\Phi} w_h w_k) + \mu (\Psi_h \Psi_k - \Psi \bar{\Psi} w_h w_k) \\ - \Omega_h w_k - \varepsilon \mu w_h \Omega_k \quad (h, k = 1, 2, 3, 4);$$

darin sind noch einzusetzen

$$\Phi \bar{\Phi} = \Phi_1^2 + \Phi_2^2 + \Phi_3^2 + \Phi_4^2, \quad \Psi \bar{\Psi} = \Psi_1^2 + \Psi_2^2 + \Psi_3^2 + \Psi_4^2, \\ e_{hh} = 1, \quad e_{hk} = 0 (h \neq k).$$

Nämlich jedenfalls ist die rechte Seite von (82) ebenso wie L eine Invariante bei den Lorentz-Transformationen und stellen die 4×4 Elemente rechts in (83) ebenso wie die S_{hk} eine Raum-Zeit-Matrix II. Art dar. Mit Rücksicht hierauf genügt es schon, um die Relationen (82), (83) allgemein behaupten zu können, sie nur für den Fall $w_1 = 0, w_2 = 0, w_3 = 0, w_4 = i$ zu verifizieren. Für diesen Fall $w = 0$ aber kommen (83) und (82) durch (47), (51), (60) einerseits, $\varepsilon = \varepsilon \mathfrak{E}$, $\mathfrak{M} = \mu \mathfrak{M}$ andererseits unmittelbar auf die Gleichungen (75) und (76) hinaus.

Der Ausdruck rechts in (81), der

$$= \left(\frac{1}{2} (\mathfrak{M} \mathfrak{M} - \varepsilon \mathfrak{E}) \right)^2 + (\varepsilon \mathfrak{M}) (\mathfrak{E} \mathfrak{M})$$

ist, erweist sich durch $(\varepsilon \mathfrak{M}) = \varepsilon \Phi \bar{\Psi}$, $(\mathfrak{E} \mathfrak{M}) = \mu \Phi \bar{\Psi}$ als ≥ 0 ; die Quadratwurzel aus ihm, ≥ 0 genommen, mag im Hinblick auf (79) mit $\text{Det}^{\frac{1}{4}} S$ bezeichnet werden.

Für $\bar{S}$, die transponierte Matrix von S , folgt aus (78), da $\bar{f} = -f$, $\bar{F} = -F$ ist,

$$(84) \quad Ff = \bar{S} - L, \quad f^* F^* = -\bar{S} - L.$$

Sodann ist

$$S - \bar{S} = |S_{hk} - S_{kh}|$$

eine alternierende Matrix und bedeutet zugleich einen Raum-Zeit-Vektor II. Art. Aus den Ausdrücken (83) entnehmen wir sofort

$$(85) \quad S - \bar{S} = -(\varepsilon\mu - 1)[w, \Omega],$$

woraus noch (vgl. (57), (58))

$$(86) \quad w(S - \bar{S})^* = 0,$$

$$(87) \quad w(S - \bar{S}) = (\varepsilon\mu - 1)\Omega$$

herzuleiten ist.

Wenn in einem Raum-Zeitpunkte die Materie ruht, $w = 0$ ist, so bedeutet (86) das Bestehen der Gleichungen

$$Z_y = Y_z, \quad X_z = Z_x, \quad Y_x = X_y;$$

ferner hat man dann nach (83):

$$\begin{aligned} T_x &= \Omega_1, & T_y &= \Omega_2, & T_z &= \Omega_3, \\ X_t &= \varepsilon\mu\Omega_1, & Y_t &= \varepsilon\mu\Omega_2, & Z_t &= \varepsilon\mu\Omega_3. \end{aligned}$$

Nun wird man durch eine geeignete Drehung des räumlichen Koordinatensystems der x, y, z um den Nullpunkt es bewirken können, daß

$$Z_y = Y_z = 0, \quad X_z = Z_x = 0, \quad Y_x = X_y = 0$$

ausfallen. Nach (71) hat man

$$(88) \quad X_x + Y_y + Z_z + T_t = 0$$

und nach dem Ausdruck in (83) ist hier jedenfalls $T_t > 0$. Im speziellen Falle, daß auch Ω verschwindet, folgt dann aus (81)

$$X_x^2 = Y_y^2 = Z_z^2 = T_t^2 = (\text{Det}^{\frac{1}{4}} S)^2$$

und sind T_t und von den drei Größen X_x, Y_y, Z_z eine $= +\text{Det}^{\frac{1}{4}} S$, die zwei anderen $= -\text{Det}^{\frac{1}{4}} S$. Verschwindet Ω nicht, so sei etwa $\Omega_3 \neq 0$, dann hat man nach (80) insbesondere

$$T_z X_t = 0, \quad T_x Y_t = 0, \quad Z_x T_z + T_z T_t = 0$$

und findet demnach $\Omega_1 = 0, \Omega_2 = 0, Z_x = -T_t$. Aus (81) und im Hinblick auf (88) folgt alsdann

$$\begin{aligned} X_x &= -Y_y = \pm \text{Det}^{\frac{1}{4}} S, \\ -Z_z &= T_t = \sqrt{\text{Det}^{\frac{1}{2}} S + \varepsilon\mu\Omega_3^2} > \text{Det}^{\frac{1}{4}} S. \end{aligned}$$

Von ganz besonderer Bedeutung wird endlich der Raum-Zeit-Vektor I. Art

$$(89) \quad K = \text{lor } S,$$

für den wir jetzt eine wichtige Umformung nachweisen wollen.

Nach (78) ist $S = L + fF$ und es folgt zunächst

$$\text{lor } S = \text{lor } L + \text{lor } fF.$$

Das Symbol lor bedeutet einen Differentiationsprozeß, der in $\text{lor } fF$ einerseits die Komponenten von f , andererseits die Komponenten von F betreffen wird. Entsprechend zerlegt sich $\text{lor } fF$ additiv in einen ersten und einen zweiten Teil. Der erste Teil wird offenbar das Produkt der Matrizen $(\text{lor } f)F$ sein, darin $\text{lor } f$ als 1×4 -reihige Matrix für sich aufgefaßt. Der zweite Teil ist derjenige Teil von $\text{lor } fF$, in dem die Differentiationen nur die Komponenten von F betreffen. Nun entnehmen wir aus (78)

$$fF = -F^*f^* - 2L;$$

infolgedessen wird dieser zweite Teil von $\text{lor } fF$ sein $-(\text{lor } F^*)f^* +$ dem Teil von $-2\text{lor } L$, in dem die Differentiationen nur die Komponenten von F betreffen. Danach entsteht

$$(90) \quad \text{lor } S = (\text{lor } f)F - (\text{lor } F^*)f^* + N,$$

wo N den Vektor mit den Komponenten

$$\begin{aligned} N_h = \frac{1}{2} \left(\frac{\partial f_{23}}{\partial x_h} F_{23} + \frac{\partial f_{31}}{\partial x_h} F_{31} + \frac{\partial f_{12}}{\partial x_h} F_{12} + \frac{\partial f_{14}}{\partial x_h} F_{14} + \frac{\partial f_{24}}{\partial x_h} F_{24} + \frac{\partial f_{34}}{\partial x_h} F_{34} \right. \\ \left. - f_{23} \frac{\partial F_{23}}{\partial x_h} - f_{31} \frac{\partial F_{31}}{\partial x_h} - f_{12} \frac{\partial F_{12}}{\partial x_h} - f_{14} \frac{\partial F_{14}}{\partial x_h} - f_{24} \frac{\partial F_{24}}{\partial x_h} - f_{34} \frac{\partial F_{34}}{\partial x_h} \right) \\ (h = 1, 2, 3, 4) \end{aligned}$$

bedeutet. Durch Benutzung der Grundgleichungen $\{A\}$ und $\{B\}$ geht (90) in die *fundamentale Relation*

$$(91) \quad \text{lor } S = -sF + N$$

über.

Im Grenzfalle $\varepsilon = 1$, $\mu = 1$, wo $f = F$ ist, verschwindet N identisch.

Allgemein gelangen wir auf Grund von (55), (56) und im Hinblick auf den Ausdruck (82) von L und auf (57) zu folgenden Ausdrücken der Komponenten von N :

$$\begin{aligned} (92) \quad N_h = -\frac{1}{2} \Phi \bar{\Phi} \frac{\partial \varepsilon}{\partial x_h} - \frac{1}{2} \Psi \bar{\Psi} \frac{\partial \mu}{\partial x_h} \\ + (\varepsilon \mu - 1) \left(\Omega_1 \frac{\partial w_1}{\partial x_h} + \Omega_2 \frac{\partial w_2}{\partial x_h} + \Omega_3 \frac{\partial w_3}{\partial x_h} + \Omega_4 \frac{\partial w_4}{\partial x_h} \right) \\ \text{für } h = 1, 2, 3, 4. \end{aligned}$$

Machen wir noch von (59) Gebrauch und bezeichnen den Raum-Vektor, der $\Omega_1, \Omega_2, \Omega_3$ als x -, y -, z -Komponenten hat, mit $\mathfrak{B}$, so kann der letzte, dritte Bestandteil von (92) auch auf die Gestalt

$$(93) \quad \frac{\varepsilon\mu - 1}{\sqrt{1 - w^2}} \left(\mathfrak{W} \frac{\partial w}{\partial x_h} \right)$$

gebracht werden, wobei die Klammer das skalare Produkt der darin aufgeführten zwei Vektoren anzeigt.

§ 14.

Die ponderomotorischen Kräfte.

Wir stellen jetzt die Relation $K = \text{lor } S = -sF + N$ ausführlicher dar; sie liefert die vier Gleichungen

$$(94) \quad K_1 = \frac{\partial X_x}{\partial x} + \frac{\partial X_y}{\partial y} + \frac{\partial X_z}{\partial z} - \frac{\partial X_t}{\partial t} = \varrho \mathfrak{E}_x + \mathfrak{s}_y \mathfrak{M}_z - \mathfrak{s}_z \mathfrak{M}_y \\ - \frac{1}{2} \Phi \bar{\Phi} \frac{\partial \varepsilon}{\partial x} - \frac{1}{2} \Psi \bar{\Psi} \frac{\partial \mu}{\partial x} + \frac{\varepsilon\mu - 1}{\sqrt{1 - w^2}} \left(\mathfrak{W} \frac{\partial w}{\partial x} \right),$$

$$(95) \quad K_2 = \frac{\partial Y_x}{\partial x} + \frac{\partial Y_y}{\partial y} + \frac{\partial Y_z}{\partial z} - \frac{\partial Y_t}{\partial t} = \varrho \mathfrak{E}_y + \mathfrak{s}_z \mathfrak{M}_x - \mathfrak{s}_x \mathfrak{M}_z \\ - \frac{1}{2} \Phi \bar{\Phi} \frac{\partial \varepsilon}{\partial y} - \frac{1}{2} \Psi \bar{\Psi} \frac{\partial \mu}{\partial y} + \frac{\varepsilon\mu - 1}{\sqrt{1 - w^2}} \left(\mathfrak{W} \frac{\partial w}{\partial y} \right),$$

$$(96) \quad K_3 = \frac{\partial Z_x}{\partial x} + \frac{\partial Z_y}{\partial y} + \frac{\partial Z_z}{\partial z} - \frac{\partial Z_t}{\partial t} = \varrho \mathfrak{E}_z + \mathfrak{s}_x \mathfrak{M}_y - \mathfrak{s}_y \mathfrak{M}_x \\ - \frac{1}{2} \Phi \bar{\Phi} \frac{\partial \varepsilon}{\partial z} - \frac{1}{2} \Psi \bar{\Psi} \frac{\partial \mu}{\partial z} + \frac{\varepsilon\mu - 1}{\sqrt{1 - w^2}} \left(\mathfrak{W} \frac{\partial w}{\partial z} \right),$$

$$(97) \quad \frac{1}{i} K_4 = -\frac{\partial T_x}{\partial x} - \frac{\partial T_y}{\partial y} - \frac{\partial T_z}{\partial z} - \frac{\partial T_t}{\partial t} = \mathfrak{s}_x \mathfrak{E}_x + \mathfrak{s}_y \mathfrak{E}_y + \mathfrak{s}_z \mathfrak{E}_z \\ + \frac{1}{2} \Phi \bar{\Phi} \frac{\partial \varepsilon}{\partial t} + \frac{1}{2} \Psi \bar{\Psi} \frac{\partial \mu}{\partial t} - \frac{\varepsilon\mu - 1}{\sqrt{1 - w^2}} \left(\mathfrak{W} \frac{\partial w}{\partial t} \right).$$

Es ist nun meine Meinung, daß bei den elektromagnetischen Vorgängen die ponderomotorische Kraft, die an der Materie in einem Raum-Zeitpunkte x, y, z, t angreift, berechnet für die Volumeneinheit, als x -, y -, z -Komponenten die drei ersten Komponenten des zum Raum-Zeit-Vektor w normalen Raum-Zeit-Vektors

$$(98) \quad K + (w \bar{K}) w$$

hat und daß ferner der Energiesatz seinen Ausdruck in der obigen vierten Relation findet.

Diese Meinung eingehend zu begründen, sei einem folgenden Aufsatz vorbehalten; hier will ich nur noch durch einige Ausführungen zur Mechanik dieser Meinung eine gewisse Stütze geben.

Im Grenzfalle $\varepsilon = 1$, $\mu = 1$, $\sigma = 0$ ist der Vektor $N = 0$, $\mathfrak{s} = \varrho w$, es wird dadurch $w \bar{K} = 0$ und es decken sich diese Ansätze mit den in der Elektronentheorie üblichen.

Anhang.

Mechanik und Relativitätspostulat.

Es wäre höchst unbefriedigend, dürfte man die neue Auffassung des Zeitbegriffs, die durch die Freiheit der Lorentz-Transformationen gekennzeichnet ist, nur für ein Teilgebiet der Physik gelten lassen.

Nun sagen viele Autoren, die klassische Mechanik stehe im Gegensatz zu dem Relativitätspostulate, das hier für die Elektrodynamik zugrunde gelegt ist.

Um hierüber ein Urteil zu gewinnen, fassen wir eine spezielle Lorentz-Transformation ins Auge, wie sie durch die Gleichungen (10), (11), (12) dargestellt ist, mit einem von Null verschiedenen Vektor $\mathbf{v}$ von irgendeiner Richtung und einem Betrage q , der < 1 ist. Wir wollen aber für einen Moment noch keine Verfügung über das Verhältnis von Längeneinheit und Zeiteinheit getroffen denken und demgemäß in jenen Gleichungen statt t, t', q schreiben $ct, ct', \frac{q}{c}$, wobei dann c eine gewisse positive Konstante vorstellt und $q < c$ sein muß. Die genannten Gleichungen verwandeln sich dadurch in

$$\mathbf{r}'_{\mathbf{v}} = \mathbf{r}_{\mathbf{v}}, \quad \mathbf{r}'_{\mathbf{v}} = \frac{c(\mathbf{r}_{\mathbf{v}} - q\mathbf{t})}{\sqrt{c^2 - q^2}}, \quad t' = \frac{-q\mathbf{r}_{\mathbf{v}} + c^2 t}{c\sqrt{c^2 - q^2}};$$

es bedeutet, wie wir erinnern, $\mathbf{r}$ den Raumvektor x, y, z und $\mathbf{r}'$ den Raumvektor x', y', z' .

Gehen wir in diesen Gleichungen, während wir $\mathbf{v}$ festhalten, zur Grenze $c = \infty$ über, so entsteht aus ihnen

$$\mathbf{r}'_{\mathbf{v}} = \mathbf{r}_{\mathbf{v}}, \quad \mathbf{r}'_{\mathbf{v}} = \mathbf{r}_{\mathbf{v}} - q\mathbf{t}, \quad t' = t.$$

Diese neuen Gleichungen würden nun bedeuten einen Übergang vom räumlichen Koordinatensysteme x, y, z zu einem anderen räumlichen Koordinatensysteme x', y', z' mit parallelen Achsen, dessen Nullpunkt in bezug auf das erste in gerader Linie mit konstanter Geschwindigkeit fortschreitet, während der Zeitparameter ganz unberührt bleiben soll.

Auf Grund dieser Bemerkung darf man sagen:

Die klassische Mechanik postuliert eine Kovarianz der physikalischen Gesetze für die Gruppe der homogenen linearen Transformationen des Ausdrucks

$$(1) \quad -x^2 - y^2 - z^2 + c^2 t^2$$

in sich mit der Bestimmung $c = \infty$.

Nun wäre es geradezu verwirrend, in einem Teilgebiet der Physik eine Kovarianz der Gesetze für die Transformationen des Ausdrucks (1)

in sich bei einem bestimmten endlichen c , in einem anderen Teilgebiete aber für $c = \infty$ zu finden. Daß die Newtonsche Mechanik nur diese Kovarianz für $c = \infty$ behaupten und sie nicht für den Fall von c als Lichtgeschwindigkeit ersinnen konnte, bedarf keiner Erklärung. Sollte aber nicht gegenwärtig der Versuch zulässig sein, jene traditionelle Kovarianz für $c = \infty$ nur als eine durch die Erfahrungen zunächst gewonnene Approximation an eine exaktere Kovarianz der Naturgesetze für ein gewisses endliches c aufzufassen?

Ich möchte ausführen, daß durch eine *Reformierung der Mechanik*, wobei an Stelle des Newtonschen Relativitätspostulates mit $c = \infty$ ein solches für ein endliches c tritt, sogar der axiomatische Aufbau der Mechanik erheblich an Vollendung zu gewinnen scheint.

Das Verhältnis der Zeiteinheit zur Längeneinheit sei derart normiert, daß das Relativitätspostulat mit $c = 1$ in Betracht kommt.

Indem ich jetzt geometrische Bilder auf die Mannigfaltigkeit der vier Variablen x, y, z, t übertragen will, mag es zum leichteren Verständnis des Folgenden bequem sein, zunächst y, z völlig außer Betracht zu lassen und x und t als irgendwelche schiefwinklige Parallelkoordinaten in einer Ebene zu deuten.

Ein Raum-Zeit-Nullpunkt O ($x, y, z, t = 0, 0, 0, 0$) wird bei den Lorentz-Transformationen festgehalten. Das Gebilde

$$(2) \quad -x^2 - y^2 - z^2 + t^2 = 1, \quad t > 0,$$

eine *hyperboloidische Schale*, umfaßt den Raum-Zeitpunkt $A(x, y, z, t = 0, 0, 0, 1)$ und alle Raum-Zeitpunkte A' , die nach Lorentz-Transformationen als $(x', y', z', t' = 0, 0, 0, 1)$ in den neu eingeführten Bestimmungsstücken x', y', z', t' auftreten.

Die Richtung eines Radiusvektors OA' von O nach einem Punkte A' von (2) und die Richtungen der in A' an (2) gehenden Tangenten sollen *normal* zueinander heißen.

Verfolgen wir eine bestimmte Stelle der Materie in ihrer Bahn zu allen Zeiten t . Die Gesamtheit der Raum-Zeitpunkte x, y, z, t , die der Stelle zu den verschiedenen Zeiten t entsprechen, nenne ich eine *Raum-Zeitlinie*.

Die Aufgabe, die Bewegung der Materie zu bestimmen, ist dahin aufzufassen: *Es soll für jeden Raum-Zeitpunkt die Richtung der daselbst durchlaufenden Raum-Zeitlinie festgestellt werden.*

Einen Raum-Zeitpunkt $P(x, y, z, t)$ auf *Ruhe transformieren*, heißt, durch eine Lorentz-Transformation ein Bezugssystem x', y', z', t' einführen derart, daß die t' -Achse OA' die Richtung erlangt, die in P die dort durchlaufende Raum-Zeitlinie zeigt. Der Raum $t' = \text{konst.}$, der durch P

zu legen ist, soll dann der in P auf der Raum-Zeitlinie *normale* Raum heißen. Dem Zuwachs dt der Zeit t von P aus entspricht der Zuwachs

$$(3) \quad d\tau = \sqrt{dt^2 - dx^2 - dy^2 - dz^2} = dt \sqrt{1 - w^2} = \frac{dx_4}{w_4}^*)$$

des hierbei einzuführenden Parameters t' . Der Wert des Integrals

$$\int d\tau = \int \sqrt{-(dx_1^2 + dx_2^2 + dx_3^2 + dx_4^2)},$$

auf der Raum-Zeitlinie von irgendeinem festen Anfangspunkte P^0 an bis zum variabel gedachten Endpunkte P gerechnet, heiße die *Eigenzeit* der betreffenden Stelle der Materie im Raum-Zeitpunkte P . (Es ist das eine Verallgemeinerung des von Lorentz für gleichförmige Bewegungen gebildeten Begriffs der *Ortszeit*.)

Nehmen wir einen räumlich ausgedehnten Körper R^0 zu einer bestimmten Zeit t^0 , so soll der Bereich aller durch die Raum-Zeitpunkte R^0, t^0 führenden Raum-Zeitlinien ein *Raum-Zeitfaden* heißen.

Haben wir einen analytischen Ausdruck $\Theta(x, y, z, t)$, so daß $\Theta(x, y, z, t) = 0$ von jeder Raum-Zeitlinie des Fadens in einem Punkte getroffen wird, wobei

$$-\left(\frac{\partial \Theta}{\partial x}\right)^2 - \left(\frac{\partial \Theta}{\partial y}\right)^2 - \left(\frac{\partial \Theta}{\partial z}\right)^2 + \left(\frac{\partial \Theta}{\partial t}\right)^2 > 0, \quad \frac{\partial \Theta}{\partial t} > 0$$

ist, so wollen wir die Gesamtheit Q der betreffenden Treffpunkte einen *Querschnitt* des Fadens nennen. An jedem Punkte $P(x, y, z, t)$ eines solchen Querschnitts können wir durch eine Lorentz-Transformation ein Bezugssystem x', y', z', t' einführen, so daß hernach

$$\frac{\partial \Theta}{\partial x'} = 0, \quad \frac{\partial \Theta}{\partial y'} = 0, \quad \frac{\partial \Theta}{\partial z'} = 0, \quad \frac{\partial \Theta}{\partial t'} > 0$$

wird. Die Richtung der betreffenden, eindeutig bestimmten t' -Achse heiße die *obere Normale* des Querschnitts Q im Punkte P und der Wert $dJ = \iiint dx' dy' dz'$ für eine Umgebung von P auf dem Querschnitt ein *Inhaltselement* des Querschnitts. In diesem Sinne ist R^0, t^0 selbst als der zur t -Achse normale Querschnitt $t = t^0$ des Fadens und das Volumen des Körpers R^0 als der *Inhalt* dieses Querschnitts zu bezeichnen.

Indem wir den Raum R^0 nach einem Punkte hin konvergieren lassen, kommen wir zum Begriffe eines *unendlich dünnen* Raum-Zeitfadens. In einem solchen denken wir uns stets eine Raum-Zeitlinie irgendwie als *Hauptlinie* ausgezeichnet und verstehen unter der *Eigenzeit des Fadens* die auf dieser Hauptlinie festgestellte Eigenzeit, unter den *Normalquerschnitten* des Fadens seine Durchquerungen durch die in den Punkten der Hauptlinie auf dieser normalen Räume.

*) Die Bezeichnung mit Indizes und die Zeichen w, w nehmen wir wieder in dem früher festgesetzten Sinne in Gebrauch (s. § 3 und § 4).

Wir formulieren nunmehr das *Prinzip von der Erhaltung der Massen*.

Jedem Raume R zu einer Zeit t gehört eine positive Größe, die *Masse in R zur Zeit t* , zu. Konvergiert R nach einem Punkte x, y, z, t hin, so nähert sich der Quotient aus dieser Masse und dem Volumen von R einem Grenzwert $\mu(x, y, z, t)$, der *Massendichte* im Raum-Zeitpunkte x, y, z, t .

Das Prinzip von der Erhaltung der Massen besagt: *Für einen unendlich dünnen Raum-Zeitfaden ist das Produkt μdJ aus der Massendichte μ an einer Stelle x, y, z, t des Fadens (d. h. der Hauptlinie des Fadens) und dem Inhalt dJ des durch die Stelle gehenden zur t -Achse normalen Querschnitts stets längs des ganzen Fadens konstant.*

Nun wird als Inhalt dJ_n des durch x, y, z, t gelegten Normalquerschnitts des Fadens

$$(4) \quad dJ_n = \frac{1}{\sqrt{1-w^2}} dJ = -iw_4 dJ = \frac{dt}{d\tau} dJ$$

zu rechnen sein und es möge

$$(5) \quad v = \frac{\mu}{-iw_4} = \mu \sqrt{1-w^2} = \mu \frac{d\tau}{dt}$$

als *Ruh-Massendichte* an der Stelle x, y, z, t definiert werden. Alsdann kann das Prinzip von der Erhaltung der Massen auch so formuliert werden:

Für einen unendlich dünnen Raum-Zeitfaden ist das Produkt aus der Ruh-Massendichte und dem Inhalt des Normalquerschnitts an einer Stelle des Fadens stets längs des ganzen Fadens konstant.

In einem beliebigen Raum-Zeitfaden sei ein erster Querschnitt Q^0 und sodann ein zweiter Querschnitt Q^1 angebracht, der mit Q^0 dessen Punkte auf der Begrenzung des Fadens, aber nur diese gemein hat, und die Raum-Zeitlinien innerhalb des Fadens mögen auf Q^1 größere Werte t als auf Q^0 zeigen. Das von Q^0 und Q^1 zusammen begrenzte, im Endlichen gelegene Gebiet soll dann eine *Raum-Zeit-Sichel*, Q^0 die *untere*, Q^1 die *obere* Begrenzung der Sichel heißen.

Denken wir uns den Faden in viele sehr dünne Raum-Zeitfäden zerlegt, so entspricht jedem Eintritt eines dünnen Fadens in die untere Begrenzung der Sichel ein Austritt aus der oberen, wobei für beide das im Sinne von (4) und (5) ermittelte Produkt $v dJ_n$ jedesmal gleichen Wert hat. Es verschwindet daher die Differenz der zwei Integrale $\int v dJ_n$, das erste erstreckt über die obere, das zweite über die untere Begrenzung der Sichel. Diese Differenz findet sich nach einem bekannten Theoreme der Integralrechnung gleich dem Integrale

$$\iiint \int \text{lor } v \bar{w} dx dy dz dt,$$

erstreckt über das ganze Gebiet der Sichel, wobei (vgl. (67) in § 12)

$$\text{lor } \nu \bar{w} = \frac{\partial \nu w_1}{\partial x_1} + \frac{\partial \nu w_2}{\partial x_2} + \frac{\partial \nu w_3}{\partial x_3} + \frac{\partial \nu w_4}{\partial x_4}$$

ist. Wird die Sichel auf einen Raum-Zeitpunkt x, y, z, t zusammengezogen, so folgt hiernach die Differentialgleichung

$$(6) \quad \text{lor } \nu \bar{w} = 0,$$

d. i. die *Kontinuitätsbedingung*

$$-\frac{\partial \mu w_x}{\partial x} + \frac{\partial \mu w_y}{\partial y} + \frac{\partial \mu w_z}{\partial z} + \frac{\partial \mu}{\partial t} = 0.$$

Wir bilden ferner, über das ganze Gebiet einer Raum-Zeit-Sichel erstreckt, das Integral

$$(7) \quad N = \iiint \nu dx dy dz dt.$$

Wir zerschneiden die Sichel in dünne Raum-Zeitfäden und jeden dieser Fäden weiter nach kleinen Elementen $d\tau$ seiner Eigenzeit, die aber noch gegen die Lineardimensionen der Normalquerschnitte groß sind, setzen die Masse eines solchen Fadens $\nu dJ_n = dm$ und schreiben noch τ^0 und τ^1 für die Eigenzeit des Fadens auf der unteren bzw. der oberen Begrenzung der Sichel; alsdann ist das Integral (7) auch zu deuten als

$$\iint \nu dJ_n d\tau = \int (\tau^1 - \tau^0) dm$$

über die sämtlichen Fäden in der Sichel.

Nun fasse ich die Raum-Zeitlinien innerhalb einer Raum-Zeit-Sichel gleichsam wie substanzielle Kurven aus substanziellen Punkten bestehend auf und denke sie mir einer kontinuierlichen Lagenveränderung innerhalb der Sichel in folgender Art unterworfen. Die ganzen Kurven sollen irgendwie *unter Festhaltung der Endpunkte auf der unteren und der oberen Begrenzung* der Sichel verrückt und die einzelnen substanziellen Punkte auf ihnen dabei so geführt werden, daß sie stets *normal zu den Kurven* fortschreiten. Der ganze Prozeß soll analytisch mittels eines Parameters ϑ darzustellen sein und dem Werte $\vartheta = 0$ sollen die Kurven in dem wirklich stattfindenden Verlauf der Raum-Zeitlinien innerhalb der Sichel entsprechen. Ein solcher Prozeß soll eine *virtuelle Verrückung in der Sichel* heißen.

Der Punkt x, y, z, t in der Sichel für $\vartheta = 0$ möge beim Parameterwerte ϑ nach $x + \delta x, y + \delta y, z + \delta z, t + \delta t$ gekommen sein; letztere Größen sind dann Funktionen von x, y, z, t, ϑ . Fassen wir wieder einen unendlich dünnen Raum-Zeitfaden an der Stelle x, y, z, t auf mit einem Normalquerschnitte von einem Inhalte dJ_n und ist $dJ_n + \delta dJ_n$ der Inhalt des Normalquerschnitts an der entsprechenden Stelle des variierten Fadens, so wollen wir dem *Prinzipie von der Erhaltung der*

Massen in der Weise Rechnung tragen, daß wir an dieser variierten Stelle eine Ruh-Massendichte $\nu + \delta\nu$ gemäß

$$(8) \quad (\nu + \delta\nu) (dJ_n + \delta dJ_n) = \nu dJ_n = dm$$

annehmen, unter ν die wirkliche Ruh-Massendichte an x, y, z, t verstanden. Zufolge dieser Festsetzung variiert dann das Integral (7), über das Gebiet der Sichel erstreckt, bei der virtuellen Verrückung als eine bestimmte Funktion $N + \delta N$ von ϑ und wir wollen diese Funktion $N + \delta N$ die *Massenwirkung* bei der virtuellen Verrückung nennen.

Ziehen wir die Schreibweise mit Indizes heran, so wird sein:

$$(9) \quad d(x_h + \delta x_h) = dx_h + \sum_k \frac{\partial \delta x_h}{\partial x_k} dx_k + \frac{\partial \delta x_h}{\partial \vartheta} d\vartheta \quad \left(\begin{matrix} k = 1, 2, 3, 4 \\ h = 1, 2, 3, 4 \end{matrix} \right).$$

Nun leuchtet auf Grund der schon gemachten Bemerkungen alsbald ein, daß der Wert von $N + \delta N$ beim Parameterwerte ϑ sein wird:

$$(10) \quad N + \delta N = \iiint \nu \frac{d(\tau + \delta\tau)}{d\tau} dx dy dz dt,$$

über die Sichel erstreckt, wobei $d(\tau + \delta\tau)$ diejenige GröÙe bedeutet, die sich aus

$$\sqrt{-(dx_1 + d\delta x_1)^2 - (dx_2 + d\delta x_2)^2 - (dx_3 + d\delta x_3)^2 - (dx_4 + d\delta x_4)^2}$$

mittels (9) und

$$dx_1 = w_1 d\tau, \quad dx_2 = w_2 d\tau, \quad dx_3 = w_3 d\tau, \quad dx_4 = w_4 d\tau, \quad d\vartheta = 0$$

ableitet; es ist also

$$(11) \quad \frac{d(\tau + \delta\tau)}{d\tau} = \sqrt{-\sum_h \left(w_h + \sum_k \frac{\partial \delta x_h}{\partial x_k} w_k \right)^2} \quad \left(\begin{matrix} k = 1, 2, 3, 4 \\ h = 1, 2, 3, 4 \end{matrix} \right).$$

Nun wollen wir den Wert des Differentialquotienten

$$(12) \quad \left(\frac{d(N + \delta N)}{d\vartheta} \right)_{(\vartheta=0)}$$

einer Umformung unterwerfen. Da jedes δx_h als Funktion der Argumente $x_1, x_2, x_3, x_4, \vartheta$ für $\vartheta = 0$ allgemein verschwindet, so ist auch allgemein

$$\frac{\partial \delta x_h}{\partial x_k} = 0 \quad \text{für } \vartheta = 0. \quad \text{Setzen wir nun}$$

$$(13) \quad \left(\frac{\partial \delta x_h}{\partial \vartheta} \right)_{\vartheta=0} = \xi_h \quad (h = 1, 2, 3, 4),$$

so folgt auf Grund von (10) und (11) für den Ausdruck (12):

$$-\iiint \nu \sum_h w_h \left(\frac{\partial \xi_h}{\partial x_1} w_1 + \frac{\partial \xi_h}{\partial x_2} w_2 + \frac{\partial \xi_h}{\partial x_3} w_3 + \frac{\partial \xi_h}{\partial x_4} w_4 \right) dx dy dz dt.$$

Für die Systeme x_1, x_2, x_3, x_4 auf der Begrenzung der Sichel sollen

$\delta x_1, \delta x_2, \delta x_3, \delta x_4$ bei jedem Werte ϑ verschwinden und sind daher auch $\xi_1, \xi_2, \xi_3, \xi_4$ überall Null. Danach verwandelt sich das letzte Integral durch partielle Integration in

$$\iiint \sum_h \xi_h \left(\frac{\partial \nu w_h w_1}{\partial x_1} + \frac{\partial \nu w_h w_2}{\partial x_2} + \frac{\partial \nu w_h w_3}{\partial x_3} + \frac{\partial \nu w_h w_4}{\partial x_4} \right) dx dy dz dt.$$

Darin ist der Klammerausdruck

$$= w_h \sum_k \frac{\partial \nu w_k}{\partial x_k} + \nu \sum_k w_k \frac{\partial w_h}{\partial x_k}.$$

Die erste Summe hier verschwindet zufolge der Kontinuitätsbedingung (6), die zweite läßt sich darstellen als

$$\frac{\partial w_h}{\partial x_1} \frac{dx_1}{d\tau} + \frac{\partial w_h}{\partial x_2} \frac{dx_2}{d\tau} + \frac{\partial w_h}{\partial x_3} \frac{dx_3}{d\tau} + \frac{\partial w_h}{\partial x_4} \frac{dx_4}{d\tau} = \frac{dw_h}{d\tau} = \frac{d}{d\tau} \left(\frac{dx_h}{d\tau} \right),$$

wobei durch $\frac{d}{d\tau}$ Differentialquotienten in Richtung der Raum-Zeitlinie einer Stelle angedeutet werden. Für den Differentialquotienten (12) resultiert damit endlich der Ausdruck

$$(14) \quad \iiint \nu \left(\frac{dw_1}{d\tau} \xi_1 + \frac{dw_2}{d\tau} \xi_2 + \frac{dw_3}{d\tau} \xi_3 + \frac{dw_4}{d\tau} \xi_4 \right) dx dy dz dt.$$

Für eine virtuelle Verrückung in der Sichel hatten wir noch die Forderung gestellt, daß die substanziell gedachten Punkte normal zu den aus ihnen hergestellten Kurven fortschreiten sollten; dies bedeutet für $\vartheta = 0$, daß die ξ_h der Bedingung

$$(15) \quad w_1 \xi_1 + w_2 \xi_2 + w_3 \xi_3 + w_4 \xi_4 = 0$$

zu entsprechen haben.

Denken wir nun an die Maxwellschen Spannungen in der Elektrodynamik ruhender Körper und betrachten wir andererseits unsere Ergebnisse in den §§ 12 und 13, so liegt eine gewisse *Anpassung des Hamiltonschen Prinzips* für kontinuierlich ausgedehnte elastische Medien an das *Relativitätspostulat* nahe.

An jedem Raum-Zeitpunkte sei (wie in § 13) eine Raum-Zeit-Matrix II. Art

$$(16) \quad S = \begin{vmatrix} S_{11} & S_{12} & S_{13} & S_{14} \\ S_{21} & S_{22} & S_{23} & S_{24} \\ S_{31} & S_{32} & S_{33} & S_{34} \\ S_{41} & S_{42} & S_{43} & S_{44} \end{vmatrix} = \begin{vmatrix} X_x & Y_x & Z_x & -iT_x \\ X_y & Y_y & Z_y & -iT_y \\ X_z & Y_z & Z_z & -iT_z \\ -iX_t & -iY_t & -iZ_t & T_t \end{vmatrix}$$

bekannt, worin $X_x, Y_x, \dots, Z_x, T_x, \dots, X_t, \dots, T_t$ reelle Größen sind.

Für eine virtuelle Verrückung in einer Raum-Zeit-Sichel bei den vorhin angewandten Bezeichnungen möge der Wert des Integrals

$$(17) \quad W + \delta W = \iiint \left(\sum_{h,k} S_{hk} \frac{\partial(x_k + \delta x_k)}{\partial x_h} \right) dx dy dz dt,$$

über das Gebiet der Sichel erstreckt, die *Spannungswirkung* bei der virtuellen Verrückung heißen.

Die hier vorkommende Summe, ausführlicher und mit reellen Größen geschrieben, ist

$$\begin{aligned} & X_x + Y_y + Z_z + T_t \\ & + X_x \frac{\partial \delta x}{\partial x} + X_y \frac{\partial \delta x}{\partial y} + \cdots + Z_z \frac{\partial \delta z}{\partial z} \\ & - X_t \frac{\partial \delta x}{\partial t} - \cdots + T_x \frac{\partial \delta t}{\partial x} + \cdots + T_t \frac{\partial \delta t}{\partial t}. \end{aligned}$$

Wir wollen nun folgendes *Minimalprinzip für die Mechanik* ansetzen:

Wird irgendeine Raum-Zeit-Sichel abgegrenzt, so soll bei jeder virtuellen Verrückung in der Sichel die Summe aus der Massenwirkung und aus der Spannungswirkung für den wirklich stattfindenden Verlauf der Raum-Zeitlinien in der Sichel stets ein Extremum sein.

Der Sinn dieser Aussage ist, daß bei jeder virtuellen Verrückung in den vorhin erklärten Zeichen

$$(18) \quad \left(\frac{d(\delta N + \delta W)}{d\delta} \right)_{\delta=0} = 0$$

sein soll.

Nach den Methoden der Variationsrechnung folgen aus diesem Minimalprinzip unter Rücksichtnahme auf die Bedingung (15) und mittels der Umformung (14) sogleich die folgenden vier Differentialgleichungen

$$(19) \quad v \frac{dw_h}{d\tau} = K_h + \kappa w_h \quad (h = 1, 2, 3, 4),$$

wo

$$(20) \quad K_h = \frac{\partial S_{1h}}{\partial x_1} + \frac{\partial S_{2h}}{\partial x_2} + \frac{\partial S_{3h}}{\partial x_3} + \frac{\partial S_{4h}}{\partial x_4}$$

die Komponenten des Raum-Zeit-Vektors I. Art $K = \text{lor } S$ sind und κ ein Faktor ist, dessen Bestimmung auf Grund von $w\bar{w} = -1$ zu erfolgen hat. Durch Multiplikation von (19) mit w_h und nachherige Summation über $h = 1, 2, 3, 4$ findet man $\kappa = K\bar{w}$ und es wird $K + (K\bar{w})w$ offenbar ein zu w normaler Raum-Zeit-Vektor I. Art. Schreiben wir die Komponenten dieses Vektors

$$X, Y, Z, iT,$$

so gelangen wir nunmehr zu folgenden *Gesetzen für die Bewegung der Materie*:

$$\begin{aligned}
 \nu \frac{d}{d\tau} \frac{dx}{d\tau} &= X, \\
 \nu \frac{d}{d\tau} \frac{dy}{d\tau} &= Y, \\
 \nu \frac{d}{d\tau} \frac{dz}{d\tau} &= Z, \\
 \nu \frac{d}{d\tau} \frac{dt}{d\tau} &= T.
 \end{aligned}
 \tag{21}$$

Dabei gilt

$$\left(\frac{dx}{d\tau}\right)^2 + \left(\frac{dy}{d\tau}\right)^2 + \left(\frac{dz}{d\tau}\right)^2 = \left(\frac{dt}{d\tau}\right)^2 - 1$$

und

$$X \frac{dx}{d\tau} + Y \frac{dy}{d\tau} + Z \frac{dz}{d\tau} = T \frac{dt}{d\tau},$$

und auf Grund dieser Umstände würde sich die vierte der Gleichungen (21) als eine Folge der drei ersten darunter ansehen lassen.

Aus (21) leiten wir weiter die Gesetze für die Bewegung eines *materiellen Punktes*, das soll heißen für den Verlauf eines unendlich dünnen Raum-Zeitfadens ab.

Es bezeichne x, y, z, t einen Punkt der im Faden irgendwie angenommenen Hauptlinie. Wir bilden die Gleichungen (21) für die Punkte des *Normalquerschnitts* des Fadens durch x, y, z, t und integrieren sie, mit dem Inhaltelement des Querschnitts multipliziert, über den ganzen Raum des Normalquerschnitts. Sind die Integrale der rechten Seiten dabei R_x, R_y, R_z, R_t und ist m die konstante Masse des Fadens, so entsteht

$$\begin{aligned}
 m \frac{d}{d\tau} \frac{dx}{d\tau} &= R_x, \\
 m \frac{d}{d\tau} \frac{dy}{d\tau} &= R_y, \\
 m \frac{d}{d\tau} \frac{dz}{d\tau} &= R_z, \\
 m \frac{d}{d\tau} \frac{dt}{d\tau} &= R_t.
 \end{aligned}
 \tag{22}$$

Dabei ist wieder R mit den Komponenten R_x, R_y, R_z, iR_t ein Raum-Zeit-Vektor I. Art, der zu dem Raum-Zeit-Vektor I. Art w , Geschwindigkeit des materiellen Punktes, mit den Komponenten

$$\frac{dx}{d\tau}, \quad \frac{dy}{d\tau}, \quad \frac{dz}{d\tau}, \quad i \frac{dt}{d\tau},$$

normal ist. Wir wollen diesen Vektor R die *bewegende Kraft* des materiellen Punktes nennen.

Integriert man jedoch die Gleichungen statt über den Normalquerschnitt des Fadens entsprechend über den zur t -Achse normalen Querschnitt des Fadens, der durch x, y, z, t gelegt ist, so entstehen (s. (4)) die Gleichungen (22), multipliziert noch mit $\frac{d\tau}{dt}$, insbesondere als letzte Gleichung darunter

$$m \frac{d}{dt} \left(\frac{dt}{d\tau} \right) = m_x R_x \frac{d\tau}{dt} + m_y R_y \frac{d\tau}{dt} + m_z R_z \frac{d\tau}{dt}.$$

Man wird nun die rechte Seite als *Arbeitsleistung* am materiellen Punkte für die Zeiteinheit aufzufassen haben. In der Gleichung selbst wird man dann den *Energiesatz* für die Bewegung des materiellen Punktes sehen und den Ausdruck

$$m \left(\frac{dt}{d\tau} - 1 \right) = m \left(\frac{1}{\sqrt{1 - v^2}} - 1 \right) = m \left(\frac{1}{2} |v|^2 + \frac{3}{8} |v|^4 + \dots \right)$$

als *kinetische Energie* des materiellen Punktes ansprechen.

Indem stets $dt \geq d\tau$ ist, könnte man den Quotienten $\frac{dt - d\tau}{d\tau}$ als das Vorgehen der Zeit gegen die Eigenzeit des materiellen Punktes bezeichnen und dann sich ausdrücken: Die kinetische Energie eines materiellen Punktes ist das Produkt seiner Masse in das Vorgehen der Zeit gegen seine Eigenzeit.

Das *Quadrupel* der Gleichungen (22) zeigt wieder die durch das Relativitätspostulat geforderte volle Symmetrie in x, y, z, it , wobei der vierten Gleichung, wie wir dies bereits in der Elektrodynamik analog antrafen, *gleichsam eine höhere physikalische Evidenz zuzuschreiben ist*. Auf Grund der Forderung dieser Symmetrie ist nach dem Muster der vierten Gleichung schon sofort das Tripel der drei ersten Gleichungen aufzubauen und im Hinblick auf diesen Umstand ist die Behauptung gerechtfertigt: *Wird das Relativitätspostulat an die Spitze der Mechanik gestellt, so folgen die vollständigen Bewegungsgesetze allein aus dem Satze von der Energie*.

Ich möchte nicht unterlassen, noch plausibel zu machen, daß nicht von den Erscheinungen der *Gravitation* her ein Widerspruch gegen die Annahme des Relativitätspostulates zu erwarten ist. *)

Ist $B^*(x^*, y^*, z^*, t^*)$ ein fester Raum-Zeitpunkt, so soll der Bereich aller derjenigen Raum-Zeitpunkte $B(x, y, z, t)$, für die

*) In einer ganz anderen Weise, als ich hier vorgehe, hat H. Poincaré (Rendiconti del Circolo Matematico di Palermo, T. XXI (1906), p. 129) das Newtonsche Attraktionsgesetz dem Relativitätspostulate anzupassen versucht.

$$(23) \quad (x - x^*)^2 + (y - y^*)^2 + (z - z^*)^2 = (t - t^*)^2, \quad t - t^* \geq 0$$

ist, das *Strahlgebilde* des Raum-Zeitpunktes B^* heißen.

Von diesem Gebilde wird eine beliebig angenommene Raum-Zeitlinie stets nur in einem einzigen Raum-Zeitpunkte B geschnitten, wie einerseits aus der *Konvexität* des Gebildes, andererseits aus dem Umstande hervorgeht, daß alle Richtungen der Raum-Zeitlinie nur Richtungen von B^* nach der konkaven Seite des Gebildes sind. Es heiße dann B^* ein *Lichtpunkt* von B .

Wird in der Bedingung (23) der Punkt $B(x, y, z, t)$ fest, der Punkt $B^*(x^*, y^*, z^*, t^*)$ variabel gedacht, so stellt die nämliche Relation den Bereich aller Raum-Zeitpunkte B^* dar, die Lichtpunkte von B sind, und es zeigt sich analog, daß auf einer beliebigen Raum-Zeitlinie stets nur ein einziger Punkt B^* vorkommt, der ein Lichtpunkt von B ist.

Es möge nun ein materieller Punkt F von der Masse m bei Vorhandensein eines anderen materiellen Punktes F^* von der Masse m^* eine bewegende Kraft nach folgendem Gesetze erfahren. Stellen wir uns die Raum-Zeitfäden von F und F^* mit Hauptlinien in ihnen vor. Es sei BC ein unendlich kleines Element der Hauptlinie von F , weiter B^* der Lichtpunkt von B , C^* der Lichtpunkt von C auf der Hauptlinie von F^* , sodann OA' der zu B^*C^* parallele Radiusvektor des hyperboloidischen Grundgebildes (2), endlich D^* der Schnittpunkt der Geraden B^*C^* mit dem durch B zu ihr normal gelegten Raume. Die bewegende Kraft des Massenpunktes F im Raum-Zeitpunkte B möge nun sein derjenige zu BC normale Raum-Zeit-Vektor I. Art, der sich additiv zusammensetzt aus dem Vektor

$$(24) \quad mm^* \left(\frac{OA'}{B^*D^*} \right)^3 BD^*$$

in Richtung BD^* und dazu einem geeigneten Vektor in Richtung B^*C^* . Dabei ist unter OA'/B^*D^* das Verhältnis der betreffenden zwei parallelen Vektoren verstanden.

Es leuchtet ein, daß diese Festsetzung einen kovarianten Charakter in bezug auf die Lorentzsche Gruppe trägt.

Wir fragen nun, wie sich hiernach der Raum-Zeitfaden von F verhält, falls der materielle Punkt F^* eine gleichförmige Translationsbewegung ausführt, d. h. die Hauptlinie des Fadens von F^* eine Gerade ist. Wir verlegen den Raum-Zeit-Nullpunkt O in sie und können durch eine Lorentz-Transformation diese Gerade als t -Achse einführen. Nun bedeute x, y, z, t den Punkt B und es sei τ^* die Eigenzeit des Punktes B^* , von O aus gerechnet. Unsere Festsetzung führt hier zu den Gleichungen

$$(25) \quad \frac{d^2 x}{d\tau^2} = -\frac{m^* x}{(t-\tau^*)^3}, \quad \frac{d^2 y}{d\tau^2} = -\frac{m^* y}{(t-\tau^*)^3}, \quad \frac{d^2 z}{d\tau^2} = -\frac{m^* z}{(t-\tau^*)^3}$$

und

$$(26) \quad \frac{d^2 t}{d\tau^2} = -\frac{m^*}{(t-\tau^*)^2} \frac{d(t-\tau^*)}{d\tau},$$

wobei

$$(27) \quad x^2 + y^2 + z^2 = (t-\tau^*)^2$$

und

$$(28) \quad \left(\frac{dx}{d\tau}\right)^2 + \left(\frac{dy}{d\tau}\right)^2 + \left(\frac{dz}{d\tau}\right)^2 = \left(\frac{dt}{d\tau}\right)^2 - 1$$

ist. Die drei Gleichungen (25) lauten in Anbetracht von (27) genau wie die Gleichungen für die Bewegung eines materiellen Punktes unter Anziehung eines festen Zentrums nach dem Newtonschen Gesetze, nur daß statt der Zeit t die Eigenzeit τ des materiellen Punktes tritt. Die vierte Gleichung (26) gibt sodann den Zusammenhang zwischen Eigenzeit und Zeit für den materiellen Punkt.

Es möge nun die Bahn des Raumpunktes x, y, z für die verschiedenen τ eine Ellipse mit der großen Halbachse a , der Exzentrizität e sein und in ihr E die exzentrische Anomalie bedeuten, τ den Zuwachs an Eigenzeit für einen vollen Umlauf in der Bahn, endlich $n\tau = 2\pi$ sein, sodaß bei geeignetem Anfangspunkte von τ die Keplersche Gleichung

$$(29) \quad n\tau = E - e \sin E$$

besteht. Verändern wir noch die Zeiteinheit und bezeichnen die Lichtgeschwindigkeit mit c , so entsteht aus (28):

$$(30) \quad \left(\frac{dt}{d\tau}\right)^2 - 1 = \frac{m^*}{ac^2} \frac{1 + e \cos E}{1 - e \cos E}.$$

Unter Vernachlässigung von c^{-4} gegen 1 folgt dann

$$ndt = nd\tau \left(1 + \frac{1}{2} \frac{m^*}{ac^2} \frac{1 + e \cos E}{1 - e \cos E}\right),$$

woraus mit Benutzung von (29) sich

$$(31) \quad nt + \text{konst.} = \left(1 + \frac{1}{2} \frac{m^*}{ac^2}\right) n\tau + \frac{m^*}{ac^2} \sin E$$

ergibt. Der Faktor $\frac{m^*}{ac^2}$ hierin ist das Quadrat des Verhältnisses einer gewissen mittleren Geschwindigkeit von F in seiner Bahn zur Lichtgeschwindigkeit. Wird für m^* die Masse der Sonne, für a die halbe große Achse der Erdbahn gesetzt, so beträgt dieser Faktor 10^{-8} .

Ein Anziehungsgesetz für Massen gemäß der eben erörterten und mit dem Relativitätspostulate verbundenen Formulierung würde zugleich eine *Fortpflanzung der Gravitation mit Lichtgeschwindigkeit* bedeuten. In Betracht der Kleinheit des periodischen Termes in (31) dürfte eine Entscheidung gegen ein solches Gesetz und die vorgeschlagene modifizierte Mechanik zugunsten des Newtonschen Attraktionsgesetzes mit der Newtonschen Mechanik aus den astronomischen Beobachtungen nicht abzuleiten sein.

Inhaltsübersicht.

	Seite
Einleitung.	
Theorie von Lorentz; Theorem, Postulat, Prinzip der Relativität	352
§ 1. Bezeichnungen	354
Erster Teil.	
Betrachtung des Grenzfalles Äther.	
§ 2. Die Grundgleichungen für den Äther	355
§ 3. Das Theorem der Relativität von Lorentz	356
§ 4. Spezielle Lorentz-Transformationen	359
§ 5. Raum-Zeit-Vektoren I. und II. Art	362
§ 6. Begriff der Zeit	365
Zweiter Teil.	
Die elektromagnetischen Vorgänge.	
§ 7. Die Grundgleichungen für ruhende Körper	366
§ 8. Die Grundgleichungen für bewegte Körper	368
§ 9. Die Grundgleichungen in der Theorie von Lorentz	372
§ 10. Die Grundgleichungen nach E. Cohn	373
§ 11. Typische Darstellung der Grundgleichungen	374
§ 12. Der Differentialoperator lor	382
§ 13. Das Produkt der Feldvektoren fF	386
§ 14. Die ponderomotorischen Kräfte	391
Anhang.	
Mechanik und Relativitätspostulat.	
Raum-Zeit-Linien, Eigenzeit, Anpassung des Hamiltonschen Prinzipes, Energiesatz und Bewegungsgleichungen, Gravitation	392

XXXI.

Eine Ableitung der Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern vom Standpunkte der Elektronentheorie.

(Aus dem Nachlaß von Hermann Minkowski bearbeitet von
Max Born in Göttingen.)

(Mathematische Annalen, Band 68, S. 526—556.)

Kurze Zeit vor seinem Tode hat mir Hermann Minkowski gesprächsweise den Grundgedanken der vorliegenden Arbeit mitgeteilt. Es handelt sich dabei um eine Ableitung der Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern, die sich nahe an den von H. A. Lorentz eingeschlagenen Weg*) anschließt; es sollen nämlich aus den im freien Äther geltenden Grundgleichungen die Gesetze für bewegte Körper dadurch hergeleitet werden, daß die Bewegungen der in der Materie eingebetteten Elektrizität (Elektronen) verfolgt werden. Minkowski behauptete damals, daß der von Lorentz eingeschlagene Weg, die Mittelwerte der von den Elektronen herrührenden Effekte zu bestimmen, mathematisch äquivalent sei einer Reihenentwicklung nach einem Parameter, der die mittlere Verschiebung der Elektronen aus ihren Ruhelagen im Inneren der Materie mißt. Einige Tage später teilte er mir mit, daß das Glied 1. Ordnung in dieser Reihe in der Tat als dielektrische Polarisierung gedeutet werden könne; seiner Überzeugung nach müsse das Glied 2. Ordnung die Magnetisierung darstellen. Aus der Beschäftigung mit diesen Gedankengängen hat ihn der Tod herausgerissen**).

*) Vgl. H. A. Lorentz, Enzyklopädie der mathematischen Wissenschaften, V 2, Art. 14, Abschnitt IV, Elektromagnetische Vorgänge in ponderablen Körpern, S. 200.

M. Abraham, Elektromagnetische Theorie der Strahlung, Leipzig 1908, 2. Aufl., 2. Abschnitt, Elektromagnetische Vorgänge in wägbaren Körpern, § 28, S. 238.

**) In seinem auf der 80. Naturforscher-Versammlung zu Köln gehaltenen Vortrage „Raum und Zeit“ (Physikalische Zeitschrift, 10. Jahrg. 1909, S. 104, und Jahresberichte

Als mir von Herrn Hilbert die auf die Elektrodynamik bezüglichen Papiere Minkowskis anvertraut wurden, habe ich sogleich gesucht, ob darin auf den genannten Gegenstand bezügliche Aufzeichnungen vorhanden seien. Ich konnte aber nur wenige Anhaltspunkte finden; denn diese über hundert eng mit Formeln bedeckten Blätter enthalten kein einziges Wort des Textes oder der Erklärung der gebrauchten Zeichen. Erst als es mir gelungen war, die Minkowskischen Ideen gemäß seinen mündlichen Mitteilungen zu rekonstruieren, habe ich in seinen Aufzeichnungen Stellen gefunden, die mit den von mir gewonnenen Formeln identisch zu sein scheinen.

Aus diesen Gründen übergebe ich diese Arbeit unter Minkowskis Namen der Öffentlichkeit. In der Ausdrucksweise und den Bezeichnungen werde ich mich nach Möglichkeit Minkowskis Gebrauche anschließen, und ich verweise dieserhalb auf seine Abhandlung: *Die Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern**).

Einleitung.

In der zitierten Abhandlung hat Minkowski die Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern mit Hilfe der in § 8, S. 368, formulierten Axiome aufgestellt. Dabei werden die in der Tat wohl nicht mehr angefochtenen Grundgleichungen für ruhende Körper (§ 7, Gleichungen (I) bis (V), S. 367) als bekannt vorausgesetzt.

Dieser Standpunkt entspricht nicht den Absichten von H. A. Lorentz, der die Vorgänge in materiellen Körpern durch geeignete Hilfshypothesen über Verschiebungen und Bewegungen der in die Materie eingelagerten Elektronen zu erklären sucht. Hierbei werden nicht die aus der Erfahrung induktiv gewonnenen Maxwell-Hertzschen Grundgleichungen für ruhende materielle Körper, sondern die für den reinen Äther von Lorentz hypothetisch angenommenen Gesetze, die eine Art Idealisierung der Maxwell'schen Gleichungen sind, als Ausgangspunkt gewählt. Diese Gesetze verknüpfen die Vektoren *elektrische Feldstärke* $\mathfrak{E}$ und *magnetische Erregung* des

der deutschen Mathematiker-Vereinigung, Bd. 18, S. 75; auch als Sonderabdruck erschienen, Leipzig, B. G. Teubner 1909; diese Ges. Abhandlungen, Bd. II, S. 431) hat Minkowski auf die Möglichkeit einer solchen elektronentheoretischen Ableitung der Grundgleichungen hingewiesen und eine Veröffentlichung darüber in Aussicht gestellt.

*) Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen, mathematisch-physikalische Klasse, 1908, S. 54; diese Ges. Abhandlungen, Bd. II, S. 352. (Im folgenden zitiert als „Grundgleichungen“; die Seitenzahlen der Zitate beziehen sich auf vorstehenden Abdruck.)

reinen Äthers $\mathfrak{M}^*$) mit dem Konvektionsstrom der Elektrizität; ist $\tilde{\varrho}$ die Raumdichte und $\mathfrak{W}$ der Raumvektor *Geschwindigkeit der Elektrizität* (der Elektronen), so lauten die Grundgleichungen für den Äther:

$$(I) \quad \text{curl } \mathfrak{M} - \frac{1}{c} \frac{\partial \mathfrak{E}}{\partial t} = \frac{1}{c} \tilde{\varrho} \mathfrak{W},$$

$$(II) \quad \text{div } \mathfrak{E} = \tilde{\varrho},$$

$$(III) \quad \text{curl } \mathfrak{E} + \frac{1}{c} \frac{\partial \mathfrak{M}}{\partial t} = 0,$$

$$(IV) \quad \text{div } \mathfrak{M} = 0.$$

Dabei ist ein geeignetes Bezugssystem rechtwinkliger Raum-Koordinaten x, y, z und der Zeit t angenommen; die Lichtgeschwindigkeit im leeren Raume ist mit c bezeichnet**).

Lorentz teilt nun die Elektronen, deren Ladungen und Bewegungen in den rechten Seiten der Gleichungen (I), (II) auftreten, in mehrere Gruppen ein. Die erste Gruppe bilden die *Leitungselektronen*, die sich wesentlich unabhängig von der Materie durch diese hindurch bewegen; sie konstituieren den „Leitungsstrom“ und ihre Ladungen bilden die „wahre Elektrizität“. Die zweite Gruppe bilden die *Polarisationselektronen*, die Gleichgewichtslagen im Inneren der materiellen Moleküle besitzen, aus denen sie durch die Einwirkung des elektromagnetischen Feldes verschoben werden können; die dadurch abgeänderte Dichte der Elektrizität wird als die der „freien Elektrizität“ bezeichnet. Die dritte Gruppe sind die *Magnetisierungselektronen*, die Umlaufbewegungen um Zentra innerhalb der Materie ausführen und, analog den Ampèreschen molekularen Kreisströmen, zu den Erscheinungen der Magnetisierung Anlaß geben. Indem nun Lorentz die Mittelwerte der Anteile des Konvektionsstromes, die von den drei Arten der Elektronen herrühren, in geeigneter Weise umformt, gelangt er zu seinen Grundgleichungen für die elektromagnetischen Vorgänge in den materiellen Körpern. Wie Minkowski***)) gezeigt hat, sind diese Gleichungen für bewegte Körper nicht mit dem Relativitätspostulat vereinbar; und die Gleichungen, die Lorentz speziell für nichtmagnetisierte Körper angibt, sind es nur approximativ, was aber seinerseits nur durch

*) In Übereinstimmung mit H. A. Lorentz haben wir die den Zustand des Äthers charakterisierenden Vektoren in dieser Weise zu bezeichnen, um sie nachher in den Grundgleichungen für materielle Körper richtig deuten zu können; daraus entspringt die Abweichung der hier gebrauchten Bezeichnung von der Minkowskis in § 2 der zitierten Arbeit.

**) Um den Vergleich mit der Lorentzschen Theorie zu erleichtern, habe ich es vermieden, $c=1$ zu setzen; durch den Gebrauch der Minkowskischen Indizesbezeichnung wird die Rechnung durch das Mitführen von c nicht belastet.

***)) Grundgleichungen, Einleitung, S. 353. Vgl. auch S. 427 vorliegender Arbeit.

zwei sich kompensierende Widersprüche gegen das Relativitätsprinzip zustande kommt*).

Indem wir uns die Aufgabe stellen, die Grundgleichungen für bewegte Körper allgemein auch unter Zulassung der Magnetisierung aus den Grundgleichungen im Äther abzuleiten, bemerken wir zuerst, daß von der charakteristischen Hypothese der Elektronentheorie, der atomistischen Struktur der Elektrizität, bei der Lorentzschen Ableitung nur ein sehr beschränkter Gebrauch gemacht wird; denn durch die Mittelwertbildung über „physikalisch unendlich kleine“ Bereiche wird diese Struktur vollständig verwischt, und die Mittelwerte, auf die es allein ankommt, werden als stetige Funktionen des Ortes und der Zeit angesehen.

Wir verzichten daher überhaupt darauf, auf die feinere Struktur der Elektrizität einzugehen. Von den Lorentzschen Vorstellungen benutzen wir nur soviel, daß wir annehmen, die *Elektrizität sei ein Kontinuum, das die Materie überall durchdringt, zum Teil sich frei innerhalb derselben bewegen kann, zum Teil aber an sie gefesselt ist und nur sehr kleine Bewegungen relativ zu ihr ausführen kann.*

Will man näheren Anschluß an Lorentz erreichen, so kann man alle im folgenden vorkommenden Größen als jene Lorentzschen Mittelwerte ansehen; es ist dann aber hier nicht nötig, sie als solche durch besondere Zeichen von den auf die einzelnen Elektronen bezogenen Größen zu unterscheiden, weil wir von den letzteren nirgends Gebrauch machen.

Um von vornherein sicher zu sein, daß alle Formeln mit dem Relativitätspostulate in Übereinstimmung sind, genügt es, bei der mathematischen Formulierung der soeben ausgesprochenen Grundannahme die von Minkowski eingeführte vierdimensionale Vektorrechnung und die von ihm statuierten Symmetrien, vor allem die der Raumkoordinaten x, y, z und der mit ci multiplizierten Zeit t , in den Vordergrund zu rücken. Dadurch wird die Kovarianz der Formeln gegenüber Lorentz-Transformationen in Evidenz gesetzt.

Ferner setzen wir voraus, daß alle vorkommenden Geschwindigkeiten kleiner als die Lichtgeschwindigkeit sind.

*) Letzteren Mangel hat Herr Ph. Frank (Annalen der Physik (4), Bd. 27, 1908, S. 1059) beseitigt, indem er, dem Lorentzschen Gedankengange sonst im wesentlichen folgend, die dem Relativitätspostulate äquivalente Kontraktionshypothese an einer früheren Stelle einführt, als es bei Lorentz geschieht.

§ 1.

Bezeichnungen.

Nach Minkowski ersetzen wir

$$x, y, z, ict$$

durch

$$x_1, x_2, x_3, x_4.$$

Mit w bezeichnen wir den Raumvektor *Geschwindigkeit der Materie* mit den Komponenten

$$w_x, w_y, w_z.$$

Aus diesen bilden wir den Raum-Zeit-Vektor I. Art w mit den Komponenten:

$$w_1 = \frac{w_x}{c \sqrt{1 - \frac{|w|^2}{c^2}}}, \quad w_2 = \frac{w_y}{c \sqrt{1 - \frac{|w|^2}{c^2}}}, \quad w_3 = \frac{w_z}{c \sqrt{1 - \frac{|w|^2}{c^2}}}, \quad w_4 = \frac{i}{\sqrt{1 - \frac{|w|^2}{c^2}}},$$

wobei $|w| = \sqrt{w_x^2 + w_y^2 + w_z^2}$ den Betrag der Geschwindigkeit bedeutet. Die Größen w_α erfüllen die Relation:

$$w_1^2 + w_2^2 + w_3^2 + w_4^2 = -1.$$

Die *Geschwindigkeit der Elektrizität* haben wir in (I) bis (IV) mit $\tilde{w}$ bezeichnet; diesem Raumvektor ordnen wir in analoger Weise einen Raum-Zeit-Vektor I. Art $\tilde{w}$ zu.

Aus der Dichte $\tilde{\varrho}$ der Elektrizität bilden wir die gegenüber Lorentz-Transformationen invariante *Ruh-Dichte*:

$$\tilde{\varrho}_0 = \tilde{\varrho} \sqrt{1 - \frac{|\tilde{w}|^2}{c^2}}.$$

Die Raumvektoren magnetische Erregung $\mathfrak{M}$ und elektrische Feldstärke $\mathfrak{E}$ fassen wir zu dem Raum-Zeit-Vektor II. Art F zusammen, indem wir

$$\mathfrak{M}_x, \mathfrak{M}_y, \mathfrak{M}_z, -i\mathfrak{E}_x, -i\mathfrak{E}_y, -i\mathfrak{E}_z$$

durch

$$F_{23}, F_{31}, F_{12}, F_{14}, F_{24}, F_{34}$$

ersetzen.

Mit Hilfe des Minkowskischen Differentialoperators lor , der als der Raum-Zeit-Vektor I. Art

$$\left| \frac{\partial}{\partial x_1}, \frac{\partial}{\partial x_2}, \frac{\partial}{\partial x_3}, \frac{\partial}{\partial x_4} \right|$$

definiert ist, können wir dann die Grundgleichungen (I) bis (IV) in folgende symbolische Gleichungen zusammenfassen:

$$(A) \quad \text{lor } F = -\tilde{\varrho}_0 \tilde{w},$$

$$(B) \quad \text{lor } F^* = 0.$$

§ 2.

Die Zerlegung der elektrischen Strömung.

Die Strömung der Elektrizität, die durch den Raum-Zeit-Vektor $\tilde{w}$ und die Ruh-Dichte $\tilde{\rho}_0$ charakterisiert ist, denken wir uns aus zwei sich superponierenden Teilen zusammengesetzt.

Der 1. Teil, der aus den Lorentzschen „Leitungselektronen“ gebildet zu denken ist, habe die Ruh-Dichte $\rho_0^{(l)}$ und der zugehörige Raum-Zeit-Vektor Geschwindigkeit sei mit $w^{(l)}$ bezeichnet. Überhaupt werden alle auf diesen Stromanteil bezüglichen Größen durch den oberen Index l gekennzeichnet. Dieser Teil des Konvektionsstromes bewege sich unabhängig von der Bewegung der Materie innerhalb derselben. Er wird zur Entstehung des Leitungsstroms Veranlassung geben.

Der 2. Teil der Strömung, der von den Lorentzschen Polarisations- und Magnetisierungselektronen gebildet ist, soll im wesentlichen der Bahn der Materie folgen und nur wenig von ihr abweichen, und zwar soll folgendes statthaben:

Von den in den materiellen neutralen Punkten vereinigten und sich dort kompensierenden Elektrizitätsmengen sei ein Quantum aus dieser Ruhelage verschoben und führe sowohl Bewegungen relativ zur Materie aus, als auch werde es von dieser mitgeführt.

Um diese Annahme zu formulieren, denken wir uns vorläufig (indem wir die das Relativitätsprinzip berücksichtigende nähere Bestimmung dieser Verschiebung für § 3 vorbehalten) sowohl die Ruh-Dichte, als auch die Komponenten des Geschwindigkeitsvektors dieses Stromanteils außer von x, y, z, t noch abhängig von einem Parameter ϑ derart, daß für $\vartheta = 0$ die Ruh-Dichte verschwindet und die Geschwindigkeitskomponenten in die entsprechenden der Materie übergehen. Wegen der vorausgesetzten Kleinheit der Abweichungen werden wir nach Potenzen von ϑ entwickeln können; bezeichnen wir die Differentiation nach ϑ bei festen x, y, z, t mit δ , so wird man für diesen zweiten Teil des Konvektionsstromes

$$\vartheta \cdot \delta(\rho_0 w) + \frac{\vartheta^2}{2} \delta \delta(\rho_0 w) + \dots$$

schreiben können. Die gesamte Strömung der Elektrizität ist demnach:

$$\tilde{\rho}_0 \tilde{w} = \rho_0^{(l)} w^{(l)} + \vartheta \delta(\rho_0 w) + \frac{\vartheta^2}{2} \delta \delta(\rho_0 w) + \dots$$

Die Glieder dieser Reihe werden wir einzeln betrachten und zeigen, daß das erste Glied, das von ϑ frei ist, den von Minkowski mit s bezeichneten Stromvektor bildet, der den Leitungsstrom und den an der Materie haftenden Konvektionsstrom der Elektrizität darstellt, daß das zweite Glied mit

dem Faktor ϑ als der Hauptteil der dielektrischen Polarisierung der Materie zu deuten ist und das dritte Glied mit ϑ^2 einerseits zu dieser Polarisierung noch einen Beitrag liefert, andererseits als Magnetisierung der Materie aufgefaßt werden muß.

Die Frage nach der Bedeutung der höheren Glieder der Reihe (1) fällt aus dem Rahmen der vorliegenden Untersuchung heraus. Bedenkt man, daß bereits die Magnetisierbarkeit bei den meisten Substanzen äußerst gering ist, derart, daß schon das in ϑ quadratische Glied der Reihe (1) einen kleinen numerischen Wert bekommt, so wird man annehmen können, daß der Einfluß der höheren Glieder auf die Beobachtungen bei den meisten Substanzen unmerklich ist, sofern ihnen überhaupt eine physikalische Bedeutung zukommt. Die Eigenschaften der stark magnetisierbaren (ferromagnetischen) Körper sind aber überhaupt noch zu wenig aufgeklärt, als daß man von ihnen aus für eine so weitgehende Theorie experimentelle Stützen erwarten dürfte.

§ 3.

Die Darstellung der variierten Strömung.

Den im vorigen Paragraphen betrachteten zweiten Teil der Strömung kann man analytisch als eine „Variation“ der Strömung der Materie ansehen*). Um diese näher zu charakterisieren, stellen wir diesen Anteil des Konvektionsstroms analog der nach Lagrange bezeichneten Weise dadurch dar, daß wir x, y, z und t als Funktionen dreier Parameter, ξ, η, ζ , welche die einzelnen materiellen Teilchen individualisieren, und der Eigenzeit τ ansehen; außerdem mögen diese vier Funktionen noch von einem Parameter ϑ derart abhängen, daß sie für $\vartheta = 0$ die Bewegung der Materie darstellen; wir schreiben also

$$(2) \quad \begin{aligned} x &= x(\xi, \eta, \zeta, \tau; \vartheta), \\ y &= y(\xi, \eta, \zeta, \tau; \vartheta), \\ z &= z(\xi, \eta, \zeta, \tau; \vartheta), \\ t &= t(\xi, \eta, \zeta, \tau; \vartheta), \end{aligned}$$

wobei identisch in allen fünf Argumenten die Bedingung

$$(3) \quad \left(\frac{\partial x}{\partial \tau}\right)^2 + \left(\frac{\partial y}{\partial \tau}\right)^2 + \left(\frac{\partial z}{\partial \tau}\right)^2 - c^2 \left(\frac{\partial t}{\partial \tau}\right)^2 = -c^2$$

erfüllt sein möge. Die Geschwindigkeitskomponenten der Materie sind

*) Diese Variation ist nicht unähnlich der virtuellen Verrückung, die Minkowski im Anhang der zitierten Arbeit S. 396 zur Ableitung der Grundgleichungen der Mechanik benützt.

$$w_1 = \frac{1}{c} \left(\frac{\partial x}{\partial \tau} \right)_{\vartheta=0}, \quad w_2 = \frac{1}{c} \left(\frac{\partial y}{\partial \tau} \right)_{\vartheta=0},$$

$$w_3 = \frac{1}{c} \left(\frac{\partial z}{\partial \tau} \right)_{\vartheta=0}, \quad w_4 = i \left(\frac{\partial t}{\partial \tau} \right)_{\vartheta=0}.$$

Ersetzen wir

$$\begin{aligned} & \xi, \eta, \zeta, i c \tau \\ \text{durch} & \xi_1, \xi_2, \xi_3, \xi_4, \end{aligned}$$

so können wir kurz schreiben:

$$(2') \quad x_\alpha = x_\alpha(\xi_1, \xi_2, \xi_3, \xi_4; \vartheta), \quad (\alpha = 1, 2, 3, 4),$$

$$(3') \quad \sum_{\alpha=1}^4 \left(\frac{\partial x_\alpha}{\partial \xi_4} \right)^2 = 1.$$

Wir wollen für die häufig vorkommenden Differentiationsprozesse folgende Abkürzungen gebrauchen. Eine Funktion φ von $x_1, x_2, x_3, x_4, \vartheta$ kann man vermöge der Transformation (2') auch als Funktion von $\xi_1, \xi_2, \xi_3, \xi_4, \vartheta$ ansehen. Wir bezeichnen

$$\frac{\partial \varphi}{\partial \xi_4} \text{ bei festgehaltenen } \xi_1, \xi_2, \xi_3, \vartheta \text{ mit } \varphi',$$

$$\frac{\partial \varphi}{\partial \vartheta} \quad " \quad " \quad \xi_1, \xi_2, \xi_3, \xi_4 \text{ mit } \dot{\varphi},$$

$$\frac{\partial \varphi}{\partial \vartheta} \quad " \quad " \quad x_1, x_2, x_3, x_4 \text{ mit } \delta \varphi.$$

Die Operationen $\varphi', \dot{\varphi}$ sind also vertauschbar; die Operation $\delta \varphi$ ist aber mit keiner der beiden ersten vertauschbar.

Die Ruh-Dichte ϱ_0 der Elektrizität wird ebenfalls eine Funktion der ξ_α und ϑ sein:

$$\varrho_0 = \varrho_0(\xi_1, \xi_2, \xi_3, \xi_4; \vartheta).$$

Da wir voraussetzen, daß die Materie vor der Verrückung, d. h. für $\vartheta = 0$, elektrisch neutral sei, müssen wir annehmen, daß

$$\varrho_0(\xi_1, \xi_2, \xi_3, \xi_4; 0) = 0$$

sei.

Dagegen setzen wir voraus, daß die Größen

$$\varrho_0 \dot{x}_1, \varrho_0 \dot{x}_2, \varrho_0 \dot{x}_3, \varrho_0 \dot{x}_4$$

sowie ihre Ableitungen nach ϑ bei festgehaltenen x_1, x_2, x_3, x_4 für $\vartheta = 0$ endliche, nicht identisch verschwindende Grenzwerte besitzen.

Daß dies mit den über die x_α und ϱ_0 gemachten Annahmen verträglich ist, zeigt z. B. der Ansatz

$$x_\alpha = f_\alpha(\xi_1, \xi_2, \xi_3, \xi_4) + \vartheta^{\frac{1}{2}} \varphi_\alpha(\xi_1, \xi_2, \xi_3, \xi_4),$$

$$\varrho_0 = \vartheta^{\frac{1}{2}} \psi(\xi_1, \xi_2, \xi_3, \xi_4);$$

hier stellt

$$x_\alpha = f_\alpha(\xi_1, \xi_2, \xi_3, \xi_4)$$

die Strömung der Materie dar, ϱ_0 verschwindet für $\vartheta = 0$; es bleibt aber

$$\varrho_0 \dot{x}_\alpha = \frac{1}{2} \varphi_\alpha \cdot \psi$$

endlich und von Null verschieden.

Der Sinn dieser Forderung ist der:

Es soll durch die Verschiebung der Ladungen aus ihrer Ruhelage jedes neutrale Teilchen der Materie in ein geladenes System (im einfachsten Falle in einen Dipol) verwandelt werden; im Sinne der Elektronentheorie werden die Größen $\varrho_0 \dot{x}_\alpha$ die „elektrischen Momente“ der Teilchen bestimmen; in der Tat werden wir diesen Umstand im folgenden klar erkennen. Daher müssen wir diesen Produkten für $\vartheta = 0$ einen endlichen Grenzwert zuschreiben. Man kann das auch so ausdrücken, daß unsere Variation der Strömung eine Art „räumlicher Doppelbelegung“ der Materie ist. Die Funktionen x_α sind infolge dieses Umstandes bei $\vartheta = 0$ nicht in Potenzreihen entwickelbar; wohl ist dieses aber, wie wir sehen werden, mit den Komponenten des Konvektionsstroms $i\varrho_0 x'_\alpha$ der Fall, auf die es allein ankommt; hier bedeuten die vier Größen $i\dot{x}_\alpha$ die Komponenten des Raum-Zeit-Vektors Geschwindigkeit; sie gehen für $\vartheta = 0$ in die Komponenten w_α der Geschwindigkeit der Materie über.

Der Unzerstörbarkeit der Elektrizität tragen wir dadurch Rechnung, daß wir die „Kontinuitätsbedingung“

$$(4) \quad \frac{\partial \varrho_0 x'}{\partial x} + \frac{\partial \varrho_0 y'}{\partial y} + \frac{\partial \varrho_0 z'}{\partial z} + \frac{\partial \varrho_0 t'}{\partial t} = 0$$

oder kurz

$$(4') \quad \sum_{\alpha=1}^4 \frac{\partial \varrho_0 x'_\alpha}{\partial x_\alpha} = 0$$

identisch in den x_α und ϑ erfüllt annehmen.

Wir müssen noch eine weitere Voraussetzung über die Größen $\dot{x}_\alpha$ machen, die die Verschiebung des elektrischen Teilchens aus seiner Ruhelage bestimmen. *Wir nehmen an, daß die Verschiebung jedes elektrischen Teilchens in einem Bezugssysteme, in dem das zu denselben Werten ξ, η, ζ gehörige materielle Teilchen ruht, durch einen Raumvektor (d. h. einen Raum-Zeit-Vektor I. Art mit verschwindender Zeitkomponente) dargestellt wird.*

Das läuft darauf heraus, daß die Raum-Zeit-Vektoren I. Art $\dot{x}$ und w normal sind (vgl. „Grundgleichungen“, § 11, 6°, S. 378):

$$(5) \quad \sum_{\alpha=1}^4 w_{\alpha} \dot{x}_{\alpha} = 0.$$

Demnach hat die Variation der Zeit $\dot{t}$ die Bedeutung des folgenden skalaren Produkts

$$(5') \quad \dot{t} = \frac{1}{c^2} (w_x \dot{x} + w_y \dot{y} + w_z \dot{z}).$$

Aus (5) ergibt sich speziell durch Multiplikation mit ϱ_0 und Differentiation nach ϑ bei festgehaltenen x, y, z, t :

$$(6) \quad \sum_{\alpha=1}^4 w_{\alpha} \delta(\varrho_0 \dot{x}_{\alpha}) = 0$$

oder

$$(6') \quad \delta(\varrho_0 \dot{t}) = \frac{1}{c^2} (w_x \delta(\varrho_0 \dot{x}) + w_y \delta(\varrho_0 \dot{y}) + w_z \delta(\varrho_0 \dot{z})).$$

Eine weitere Normalitätsbedingung erhält man, wenn man die Identität (3) bzw. (3') nach ϑ bei festen x, y, z, t differenziert:

$$(7) \quad \sum_{\alpha=1}^4 x'_{\alpha} \delta x'_{\alpha} = 0.$$

Geht man hier zur Grenze $\vartheta = 0$ über, so kommt:

$$(7') \quad \sum_{\alpha=1}^4 w_{\alpha} \delta w_{\alpha} = 0.$$

Dabei haben die δw_{α} die folgende Bedeutung:

$$(8) \quad \begin{aligned} \delta w_1 &= \delta \left(\frac{w_x}{c \sqrt{1 - \frac{|w|^2}{c^2}}} \right) = i[\delta x'_1]_{\vartheta=0}, \text{ u. s. f.} \\ \delta w_4 &= \delta \left(\frac{i}{\sqrt{1 - \frac{|w|^2}{c^2}}} \right) = i[\delta x'_4]_{\vartheta=0}. \end{aligned}$$

Endlich betrachten wir die *Variation der Ruh-Dichte* ϱ_0 . Diese soll nicht unabhängig variieren, sondern so, daß ihre Ableitung nach ϑ an einer Stelle x, y, z, t mit dem elektrischen Momente $\varrho_0 \dot{x}$ an dieser Stelle durch die Identität in den x_{α} und ϑ

$$(9) \quad \frac{\partial \varrho_0 \dot{x}}{\partial x} + \frac{\partial \varrho_0 \dot{y}}{\partial y} + \frac{\partial \varrho_0 \dot{z}}{\partial z} + \frac{\partial \varrho_0 \dot{t}}{\partial t} = -\delta \varrho_0$$

oder

$$(9') \quad \sum_{\alpha=1}^4 \frac{\partial \varrho_0 \dot{x}_{\alpha}}{\partial x_{\alpha}} = -\delta \varrho_0$$

verbunden ist.

Die Bedeutung dieser Gleichung ist die folgende: Wir nahmen an, daß das verschobene Quantum Elektrizität vor der Verschiebung durch die gleiche Menge entgegengesetzter Elektrizität kompensiert wurde; daher ist die linke Seite von Gleichung (9) nicht Null, was bedeuten würde, daß die Abnahme der in einem Volumen befindlichen Elektrizitätsmenge bei der Verschiebung gleich der durch die Begrenzung des Volumens tretenden Menge ist; vielmehr tritt bei der Verschiebung jene vorher kompensierte Ladung des entgegengesetzten Vorzeichens zutage, und das ist eben in erster Näherung die bei festgehaltenen x, y, z, t genommene Ableitung $-\delta\varphi_0$. Die spezielle Form des obigen Ansatzes wird durch die vom Relativitätsprinzip geforderte Symmetrie gerechtfertigt.

§ 4.

Die Reihenentwicklung der variierten Strömung.

Unsere Aufgabe ist es, die Größen $\varphi_0 x'_\alpha$ in jedem Raum-Zeitpunkte, d. h. bei festgehaltenen x_α , in einer Potenzreihe nach ϑ zu entwickeln; gemäß der Bedeutung des Symbols δ haben wir also:

$$(10) \quad \varphi_0 x'_\alpha = \vartheta \delta(\varphi_0 x'_\alpha)_0 + \frac{\vartheta^2}{2} \delta \delta(\varphi_0 x'_\alpha)_0 + \dots,$$

wo in den Koeffizienten $\vartheta = 0$ zu setzen ist.

Es ist

$$\delta(\varphi_0 x'_\alpha) = \varphi_0 \delta x'_\alpha + x'_\alpha \delta \varphi_0.$$

Wir wollen die Operation δ durch den mit dem Punkte bezeichneten „substantiellen“ Differentialquotienten ersetzen. Dazu sehen wir x'_α als Funktion der x_α und ϑ an, wobei die x_α ihrerseits Funktionen der ξ_α und ϑ sind:

$$x'_\alpha = x'_\alpha(x_1(\xi_1, \xi_2, \xi_3, \xi_4; \vartheta), \dots, x_4(\xi_1, \xi_2, \xi_3, \xi_4; \vartheta); \vartheta).$$

Dann ergibt sich durch Differentiation nach ϑ bei festen ξ_α :

$$(\alpha) \quad \dot{x}'_\alpha = \sum_{\beta=1}^4 \frac{\partial x'_\alpha}{\partial x_\beta} \dot{x}_\beta + \delta x'_\alpha.$$

Andererseits bilden wir dieselbe Größe $\dot{x}'_\alpha$, indem wir zuerst die mit dem Punkte, dann die mit dem Striche bezeichnete Operation vornehmen; sehen wir also $\dot{x}_\alpha$ als Funktion der x_α und ϑ an und die x_α ihrerseits als Funktionen der ξ_α und ϑ :

$$\dot{x}_\alpha = \dot{x}_\alpha(x_1(\xi_1, \xi_2, \xi_3, \xi_4; \vartheta), \dots, x_4(\xi_1, \xi_2, \xi_3, \xi_4; \vartheta); \vartheta),$$

so folgt durch Differentiation nach ξ_4 bei festen $\xi_1, \xi_2, \xi_3, \vartheta$:

$$(\beta) \quad \dot{x}'_\alpha = \sum_{\beta=1}^4 \frac{\partial \dot{x}_\alpha}{\partial x_\beta} x'_\beta.$$

Aus den Gleichungen (α) , (β) finden wir:

$$(\gamma) \quad \delta x'_\alpha = \sum_{\beta=1}^4 \left(\frac{\partial \dot{x}_\alpha}{\partial x_\beta} x'_\beta - \frac{\partial x'_\alpha}{\partial x_\beta} \dot{x}_\beta \right).$$

Um $\delta(\varrho_0 x'_\alpha)$ zu erhalten, haben wir dies mit ϱ_0 zu multiplizieren und die mit x'_α multiplizierte Gleichung (9') hinzuzufügen; wir addieren außerdem noch die mit $\dot{x}_\alpha$ multiplizierte Kontinuitätsgleichung (4'), so daß wir schließlich erhalten:

$$\delta(\varrho_0 x'_\alpha) = \sum_{\beta=1}^4 \left(\frac{\partial \dot{x}_\alpha}{\partial x_\beta} \varrho_0 x'_\beta - \frac{\partial x'_\alpha}{\partial x_\beta} \varrho_0 \dot{x}_\beta - \frac{\partial \varrho_0 \dot{x}_\beta}{\partial x_\beta} x'_\alpha + \frac{\partial \varrho_0 x'_\beta}{\partial x_\beta} \dot{x}_\alpha \right)$$

oder

$$(11) \quad \delta(\varrho_0 x'_\alpha) = \sum_{\beta=1}^4 \frac{\partial}{\partial x_\beta} (\varrho_0 \dot{x}_\alpha \cdot x'_\beta - \varrho_0 \dot{x}_\beta \cdot x'_\alpha).$$

Diese Gleichung gilt identisch in den x_α und ϑ . Wenn wir sie nochmals nach ϑ bei festen x_α differenzieren, können wir diese Operation unter den in der Summe vorkommenden Differentiationszeichen nach x_β ausführen. Daher wird:

$$(12) \quad \delta \delta(\varrho_0 x'_\alpha) = \sum_{\beta=1}^4 \frac{\partial}{\partial x_\beta} \left\{ [\delta(\varrho_0 \dot{x}_\alpha) \cdot x'_\beta - \delta(\varrho_0 \dot{x}_\beta) \cdot x'_\alpha] + [\varrho_0 \dot{x}_\alpha \delta x'_\beta - \varrho_0 \dot{x}_\beta \delta x'_\alpha] \right\}.$$

Jetzt können wir die Reihe (10) in folgende Gestalt setzen:

$$(13) \quad \varrho_0 x'_\alpha = \sum_{\beta=1}^4 \frac{\partial}{\partial x_\beta} \left\{ \left[\vartheta \varrho_0 \dot{x}_\alpha + \frac{\vartheta^2}{2} \delta(\varrho_0 \dot{x}_\alpha) \right] x'_\beta - \left[\vartheta \varrho_0 \dot{x}_\beta + \frac{\vartheta^2}{2} \delta(\varrho_0 \dot{x}_\beta) \right] x'_\alpha \right\} \\ + \sum_{\beta=1}^4 \frac{\partial}{\partial x_\beta} \left\{ \frac{\vartheta^2}{2} \varrho_0 \dot{x}_\alpha \delta x'_\beta - \frac{\vartheta^2}{2} \varrho_0 \dot{x}_\beta \delta x'_\alpha \right\} + \dots,$$

wobei in den Faktoren von ϑ , ϑ^2 , ... überall $\vartheta = 0$ zu setzen ist; dabei haben die Größen $\varrho_0 \dot{x}_\alpha$, $\delta(\varrho_0 \dot{x}_\alpha)$ endliche, nicht identisch verschwindende Grenzwerte. Gemäß Verabredung brechen wir die Reihe mit den hingeschriebenen Gliedern ab.

§ 5.

Formale Herstellung der in der bewegten Materie gültigen Differentialgleichungen.

Wir setzen zur Abkürzung:

$$(14) \quad \vartheta(\varrho_0 \dot{x}_\alpha)_{\vartheta=0} + \frac{\vartheta^2}{2} \delta(\varrho_0 \dot{x}_\alpha)_{\vartheta=0} = p_\alpha,$$

ferner, indem wir uns erinnern, daß

$$\text{ist:} \quad i(x'_\alpha)_{\mathfrak{F}=0} = w_\alpha, \quad i(\delta x'_\alpha)_{\mathfrak{F}=0} = \delta w_\alpha$$

$$(15) \quad w_\alpha p_\beta - w_\beta p_\alpha = P_{\alpha\beta},$$

$$(16) \quad \frac{\mathfrak{P}^2}{2} \{ \delta w_\alpha (q_0 \dot{x}_\beta)_0 - \delta w_\beta (q_0 \dot{x}_\alpha)_0 \} = Q_{\alpha\beta}.$$

Diese Größen sind die Komponenten der aus den Vektoren erster Art w, p bzw. $\delta w, \frac{\mathfrak{P}^2}{2} (q_0 \dot{x})_0$ durch „vektorielle“ Multiplikation gebildeten Raum-Zeit-Vektoren II. Art:

$$(15') \quad [w, p] = P,$$

$$(16') \quad \frac{\mathfrak{P}^2}{2} [\delta w, (q_0 \dot{x})_0] = Q.$$

Jetzt fügen wir den Ausdruck (10) zu $q_0^{(i)} w_\alpha^{(i)}$ hinzu und können, wenn wir noch

$$(17) \quad q_0^{(i)} w_\alpha^{(i)} = s_\alpha$$

setzen, die Komponenten des gesamten Konvektionsstromes (1) in folgender Weise schreiben:

$$(18) \quad \tilde{q}_0 \tilde{w}_\alpha = s_\alpha - \sum_{\beta=1}^4 \frac{\partial}{\partial x_\beta} (P_{\alpha\beta} + Q_{\alpha\beta}),$$

oder in der Minkowskischen Symbolik:

$$(18') \quad \tilde{q}_0 \tilde{w} = s + \text{lor} (P + Q).$$

Setzen wir das in die Grundgleichung (A) ein, so geht diese über in

$$(19) \quad \text{lor} (F + P + Q) = -s.$$

Führen wir den Raum-Zeit-Vektor II. Art

$$(20) \quad f = F + P + Q$$

ein, so können wir den Grundgleichungen die Gestalt geben:

$$\{A\} \quad \text{lor } f = -s,$$

$$\{B\} \quad \text{lor } F^* = 0.$$

Damit haben wir formal die in der bewegten Materie geltenden Differentialgleichungen gewonnen (siehe Grundgleichungen, § 12, Formeln {A} {B}, Seite 385).

Ersetzen wir

$$\text{durch} \quad f_{23}, f_{31}, f_{12}, f_{14}, f_{24}, f_{34}$$

$$m_x, m_y, m_z, -ie_x, -ie_y, -ie_z,$$

und

durch

$$s_1, \quad s_2, \quad s_3, \quad s_4$$

$$\frac{1}{c} \mathfrak{s}_x, \quad \frac{1}{c} \mathfrak{s}_y, \quad \frac{1}{c} \mathfrak{s}_z, \quad i\rho$$

und fassen die hier vorkommenden Größen bzw. zu den Raumvektoren $\mathfrak{m}$, $\mathfrak{e}$, $\mathfrak{s}$ zusammen, so können wir den Gleichungen $\{A\}$, $\{B\}$ die reelle Gestalt geben:

$$(I') \quad \text{curl } \mathfrak{m} - \frac{1}{c} \frac{\partial \mathfrak{e}}{\partial t} = \frac{1}{c} \mathfrak{s},$$

$$(II') \quad \text{div } \mathfrak{e} = \rho,$$

$$(III') \quad \text{curl } \mathfrak{E} + \frac{1}{c} \frac{\partial \mathfrak{M}}{\partial t} = 0,$$

$$(IV') \quad \text{div } \mathfrak{M} = 0.$$

Das sind die in ruhenden oder bewegten Körpern geltenden Differentialgleichungen, wenn die Vektoren $\mathfrak{m}$, $\mathfrak{e}$, $\mathfrak{s}$ bzw. als die magnetische Feldstärke, die dielektrische Erregung, der Strom und die Größe ρ als die Dichte der Elektrizität gedeutet werden dürfen. Es wird unsere Aufgabe sein, aus den Definitionsgleichungen (15), (16), (17), (20) diese Bedeutung herauszulesen; die letzteren Gleichungen müssen ferner die Beziehungen enthalten, welche die beiden Vektoren „Erregung“ mit den beiden Vektoren „Feldstärke“ verbinden und in welche die Materialeigenschaften eingehen. Allerdings werden wir sehen, daß, genau wie bei Lorentz, diese Beziehungen hier viel allgemeiner sind als die von Minkowski behandelten, der sich auf isotrope, dispersionsfreie Körper beschränkte; unsere Gleichungen können durch geeignete Zusatzhypothesen den mannigfaltigen Eigenschaften der Substanzen angepaßt werden, und die einfachste Annahme führt auf die von Minkowski erhaltenen Formeln. Um diese Sachlage zu überblicken, behandeln wir zunächst den einfachen Fall ruhender Körper.

§ 6.

Ruhende Körper.

Ruht die Materie, ist also $w = 0$, so hat nach § 1 der Raum-Zeit-Vektor w die Komponenten

$$(21) \quad w_1 = 0, \quad w_2 = 0, \quad w_3 = 0, \quad w_4 = i.$$

Der Vektor $\mathfrak{s}$ hat hier ohne weiteres die Bedeutung der *Stromstärke* des frei durch die Materie fließenden *Leitungsstroms*; die von diesem transportierte Elektrizitätsmenge $\rho = \rho^{(v)}$ heißt gewöhnlich „wahre Elektrizität“, während sie von Minkowski kurzweg *elektrische Ladung* genannt wird.

Aus den Orthogonalitätsbedingungen (5), (5') und (6), (6') folgt, daß für $w = 0$

$$(22) \quad p_4 = 0$$

ist. Die Bedeutung der drei ersten Komponenten von p ,

$$p_1 = \mathfrak{p}_x, \quad p_2 = \mathfrak{p}_y, \quad p_3 = \mathfrak{p}_z,$$

ist leicht zu erkennen. So ist z. B. das erste Glied $\vartheta(\varrho_0 \dot{x}_\alpha)_{\vartheta=0}$ von p_α in erster Annäherung das elektrische Moment eines Teilchens der Materie, und das zweite Glied von p_α liefert die für jeden Raum-Zeitpunkt x, y, z, t dazu tretende Korrektur, wenn man die Annäherung einen Schritt weiter treibt. Demnach sind die Größen $\mathfrak{p}_x, \mathfrak{p}_y, \mathfrak{p}_z$ als die Komponenten des Vektors $\mathfrak{p}$, *dielektrische Polarisation*, aufzufassen.

Die Komponenten des Raum-Zeit-Vektors II. Art P reduzieren sich für $w = 0$ wegen (21) und (22) auf:

$$(23) \quad P_{23} = 0, \quad P_{31} = 0, \quad P_{12} = 0, \quad P_{14} = -i\mathfrak{p}_x, \quad P_{24} = -i\mathfrak{p}_y, \quad P_{34} = -i\mathfrak{p}_z.$$

Infolge der Normalitätsbedingungen (5) und (7) sind für $w = 0$

$$(24) \quad (\varrho_0 \dot{x}_\alpha)_{\vartheta=0} = 0, \quad \delta w_4 = 0.$$

Daher werden die Komponenten des Raum-Zeit-Vektors II. Art Q :

$$(25) \quad \begin{aligned} Q_{23} &= \frac{1}{c} \frac{\vartheta^2}{2} (\delta w_y (\varrho_0 \dot{z})_0 - \delta w_z (\varrho_0 \dot{y})_0), \quad Q_{14} = 0, \\ Q_{31} &= \frac{1}{c} \frac{\vartheta^2}{2} (\delta w_x (\varrho_0 \dot{z})_0 - \delta w_z (\varrho_0 \dot{x})_0), \quad Q_{24} = 0, \\ Q_{12} &= \frac{1}{c} \frac{\vartheta^2}{2} (\delta w_x (\varrho_0 \dot{y})_0 - \delta w_y (\varrho_0 \dot{x})_0), \quad Q_{34} = 0. \end{aligned}$$

Die ersten drei dieser Größen sind das mit $\frac{1}{c}$ multiplizierte Vektorprodukt der Relativgeschwindigkeit der Ladung gegen die Materie in das elektrische Moment; bezeichnet man den letzteren Vektor, dessen Komponenten $\vartheta(\varrho_0 \dot{x})_0, \vartheta(\varrho_0 \dot{y})_0, \vartheta(\varrho_0 \dot{z})_0$ sind, mit $\varrho \dot{\mathfrak{r}}$ und setzt man

$$(26) \quad \mathfrak{q} = \frac{1}{2c} [\varrho \dot{\mathfrak{r}}, \delta w],$$

so wird dieser Vektor als *magnetisches Moment**) oder *Magnetisierung* zu bezeichnen sein, und man hat:

$$(25') \quad Q_{23} = -q_x, \quad Q_{31} = -q_y, \quad Q_{12} = -q_z.$$

Jetzt können wir die Gleichung (20) in folgende zwei Vektorgleichungen zerlegen:

$$(27) \quad \begin{aligned} \mathfrak{m} &= \mathfrak{M} - \mathfrak{q}, \\ \mathfrak{e} &= \mathfrak{E} + \mathfrak{p}. \end{aligned}$$

*) Vgl. z. B. M. Abraham, Elektromagnetische Theorie der Strahlung, 2. Aufl. 1908, § 28, Formel (135), S. 248.

Das ist der von Lorentz aufgestellte Zusammenhang zwischen den Vektoren m , e einerseits und $\mathfrak{M}$, $\mathfrak{E}$ andererseits*); die Polarisierung p und die Magnetisierung q sind dabei als Vektoren anzusehen, die den durch das elektromagnetische Feld hervorgebrachten Zustand der Materie charakterisieren und demnach von der Art dieser Materie abhängen. Dadurch, daß man die Freiheit hat, diese Abhängigkeit näher zu bestimmen, sei es durch elektronentheoretische Betrachtungen, sei es durch hypothetische Ansätze, kann man die mannigfachen Eigenschaften der Materie dem System der Elektrodynamik einordnen.

Wir wollen uns auf die einfachsten *Zusatzhypothesen* beschränken, die zur Beschreibung *isotroper, nichtdispersierender Körper* dienen:

1. Die Geschwindigkeit des frei beweglichen Teils des Konvektionsstromes, $w^{(1)}$, ist in jedem Raum-Zeitpunkte proportional der elektrischen Feldstärke; daraus folgt sofort das Ohmsche Gesetz:

$$(28) \quad \mathfrak{s} = \sigma \mathfrak{E},$$

wo σ die Leitfähigkeit bedeutet.

2. Die Verschiebung der an der Materie haftenden Ladungen aus ihrer Ruhelage ist proportional der elektrischen Feldstärke; daraus ergibt sich

$$(29) \quad p = (\epsilon - 1)\mathfrak{E}, \quad e = \epsilon \mathfrak{E},$$

wo ϵ die Dielektrizitätskonstante ist.

3. Das Drehmoment der an der Materie haftenden Ladungen um ihre Ruhelage ist proportional der magnetischen Feldstärke m ; demnach ist

$$(30) \quad q = (\mu - 1)m, \quad \mathfrak{M} = \mu m,$$

wo μ die magnetische Permeabilität bedeutet.

σ , ϵ und μ können Funktionen von x , y , z , t sein.

Während die Hypothesen 1 und 2 eine elektronentheoretische Deutung erlauben, ist das bei der dritten Hypothese noch nicht einwandfrei durchgeführt worden. Verallgemeinert man die Hypothese 2 derart, daß man den Elektronen Trägheit zuschreibt, so gelangt man zur Dispersionstheorie.

§ 7.

Der Leitungsstrom in bewegten Körpern.

Bewegt sich die Materie, so wird man den Vektor $\mathfrak{s}$ nicht mit dem Leitungsstrom identifizieren; vielmehr wird es nur auf die Relativbewegung

*) Vgl. H. A. Lorentz, Enzyklopädie der mathematischen Wissenschaften, V 2, Art. 14, Formeln XXX, XXV und XXI, S. 208, 209.

der Elektrizität gegen die Materie ankommen, so daß der Vektor

$$(31) \quad \mathfrak{S} = \mathfrak{s} - q\mathfrak{w} = q^{(l)}(\mathfrak{w}^{(l)} - \mathfrak{w})$$

als *Leitungsstrom* zu bezeichnen ist.

Will man jetzt durch eine Zusatzhypothese den Leitungsstrom mit der wirkenden elektrischen Feldstärke in Verbindung setzen, so entspricht es nicht dem Relativitätsprinzip, die im ruhenden Koordinatensystem gemessene elektrische Feldstärke $\mathfrak{E}$ in Zusammenhang zu bringen mit der im ruhenden Koordinatensystem gemessenen Relativgeschwindigkeit $\mathfrak{w}^{(l)} - \mathfrak{w}$. Vielmehr wirkt auf die bewegte Ladung die *elektrische Ruh-Kraft*

$$\Phi = -wF$$

(vgl. Grundgleichungen, § 11, Formeln (47), (48), Seite 380); ihre ersten drei Komponenten sind die x -, y -, z -Komponenten des Raumvektors

$$\frac{\mathfrak{E} + \frac{1}{c}[\mathfrak{w}\mathfrak{M}]}{\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}},$$

dessen Zähler auch nach Lorentz die im Äther auf *bewegte* Ladungen wirkende Kraft darstellt, ihre vierte Komponente ist

$$\Phi_4 = \frac{i(\mathfrak{w}\mathfrak{E})}{c\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}},$$

und der Vektor Φ ist normal zu w :

$$w\bar{\Phi} = w_1\Phi_1 + w_2\Phi_2 + w_3\Phi_3 + w_4\Phi_4 = 0.$$

Andererseits werden wir als „Relativgeschwindigkeit“ im Sinne des Relativitätsprinzips den zu dem Vektor w *normalen* Raum-Zeit-Vektor I. Art

$$v = w^{(l)} + (w\bar{w}^{(l)})w$$

bezeichnen.

Sehen wir v_1, v_2, v_3 als Komponenten des Raumvektors

$$\frac{\mathfrak{v}}{c\sqrt{1 - \frac{|\mathfrak{w}^{(l)}|^2}{c^2}}}$$

an, so findet man leicht die Komponenten des Raumvektors $\mathfrak{v}$ in der Richtung von $\mathfrak{w}$ und in einer zu $\mathfrak{w}$ senkrechten Richtung $\bar{\mathfrak{w}}$:

$$\mathfrak{v}_{\mathfrak{w}} = \frac{\mathfrak{w}_{\mathfrak{w}}^{(l)} - |\mathfrak{w}|}{1 - \frac{|\mathfrak{w}|^2}{c^2}},$$

$$\mathfrak{v}_{\bar{\mathfrak{w}}} = \mathfrak{w}_{\bar{\mathfrak{w}}}^{(l)}.$$

Daher hängt der Vektor $\mathfrak{v}$ mit dem Leitungsstrom $\mathfrak{J}$ folgendermaßen zusammen:

$$\varrho^{(l)} \mathfrak{v}_w = \frac{\mathfrak{J}_w}{1 - \frac{|\mathfrak{w}|^2}{c^2}},$$

$$\varrho^{(l)} \mathfrak{v}_{\bar{w}} = \mathfrak{J}_{\bar{w}}.$$

Dann gelangen wir zu dem Ohmschen Gesetze für bewegte Körper durch die *Zusatzhypothese*:

Die Relativgeschwindigkeit des frei beweglichen Konvektionsstromes ist proportional der elektrischen Ruh-Kraft; daraus folgt:

$$\varrho^{(l)} v = - \frac{\sigma}{c} w F,$$

oder

$$\{E\} \quad s + (w\bar{s})w = - \frac{\sigma}{c} w F.$$

Das ist die von Minkowski aufgestellte Relation $\{E\}$, § 12, Seite 385; sie ist äquivalent mit den Vektorgleichungen:

$$(E) \quad \frac{\mathfrak{J}_w - |\mathfrak{w}| \varrho}{\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}} = \sigma \left(\mathfrak{E} + \frac{1}{c} [\mathfrak{w} \mathfrak{M}] \right)_w,$$

$$\mathfrak{J}_{\bar{w}} = \frac{\sigma \left(\mathfrak{E} + \frac{1}{c} [\mathfrak{w} \mathfrak{M}] \right)_{\bar{w}}}{\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}.$$

§ 8.

Die dielektrische Polarisation in bewegten Körpern.

Aus der Gleichung (15')

$$P = [w, p]$$

folgt:

$$w P = w [w, p] = - (w \bar{p}) w + (w \bar{w}) p.$$

Nun erkennt man, daß infolge der Gleichungen (5) und (6) der durch (14) definierte Vektor p normal zu w ist, d. h. daß

$$(32) \quad w \bar{p} = 0$$

ist; da außerdem

$$w \bar{w} = -1$$

ist, bekommt man:

$$(33) \quad w P = -p.$$

Der Vektor I. Art p steht also zu dem Vektor II. Art P in genau derselben Beziehung wie die elektrische Ruh-Kraft Φ zu dem Feldvektor II. Art F . Wir werden daher p gemäß seiner Definition durch die Gleichungen (14) als „*dielektrische Ruh-Polarisation*“ und P als „*dielektrische Polarisation*“ bezeichnen; letztere ist also ein Vektor II. Art im Gegensatz zu der gewöhnlichen Auffassung.

Setzen wir wie im vorigen Paragraphen

$$p_1 = p_x, \quad p_2 = p_y, \quad p_3 = p_z,$$

so wird infolge von (32):

$$(32') \quad p_4 = \frac{i}{c} (w_x p_x + w_y p_y + w_z p_z) = \frac{i}{c} (w p).$$

Dann sind die Komponenten von P :

$$(34) \quad \begin{aligned} P_{23} &= \frac{w_y p_z - w_z p_y}{c \sqrt{1 - \frac{|w|^2}{c^2}}}, \quad P_{31} = \frac{w_z p_x - w_x p_z}{c \sqrt{1 - \frac{|w|^2}{c^2}}}, \quad P_{12} = \frac{w_x p_y - w_y p_x}{c \sqrt{1 - \frac{|w|^2}{c^2}}}, \\ P_{14} &= -i \frac{p_x - \frac{w_x}{c^2} (w p)}{\sqrt{1 - \frac{|w|^2}{c^2}}}, \quad P_{24} = -i \frac{p_y - \frac{w_y}{c^2} (w p)}{\sqrt{1 - \frac{|w|^2}{c^2}}}, \quad P_{34} = -i \frac{p_z - \frac{w_z}{c^2} (w p)}{\sqrt{1 - \frac{|w|^2}{c^2}}}. \end{aligned}$$

Sie setzen sich also aus den Komponenten der Raumvektoren

$$(34') \quad \mathfrak{R} = \frac{[w p]}{c \sqrt{1 - \frac{|w|^2}{c^2}}}, \quad \mathfrak{P} = \frac{p - \frac{w}{c^2} (w p)}{\sqrt{1 - \frac{|w|^2}{c^2}}}$$

in genau derselben Weise zusammen wie die Komponenten von F aus denen von $\mathfrak{M}$ und $\mathfrak{E}$. Hier können wir $\mathfrak{R}$ als *Röntgen-Vektor* bezeichnen; denn er gibt Anlaß zur Entstehung des sog. Röntgenstroms. In der Tat stimmt, wie wir in § 10 sehen werden, der Vektor $\mathfrak{R}$ bis auf Glieder zweiter Ordnung in $\frac{w}{c}$ überein mit demjenigen Vektor, den auch Lorentz zur Erklärung der von Röntgen entdeckten magnetischen Wirkung polarisierter, bewegter Dielektrika erhält. Der Vektor $\mathfrak{P}$, den man *Raumvektor Polarisation* nennen könnte, unterscheidet sich von der Ruh-Polarisation p nur durch Größen 2. Ordnung. Die Komponenten von $\mathfrak{P}$ in der Richtung w und einer auf w senkrechten Richtung $\bar{w}$ hängen mit den entsprechenden Komponenten von p so zusammen:

$$\mathfrak{P}_w = p_w \sqrt{1 - \frac{|w|^2}{c^2}}, \quad \mathfrak{P}_{\bar{w}} = \frac{p_{\bar{w}}}{\sqrt{1 - \frac{|w|^2}{c^2}}}.$$

Die für isotrope, nichtdispargierende Körper gültige *Zusatzhypothese* ist hier so zu formulieren:

Der Raum-Zeit-Vektor I. Art Ruh-Polarisation p ist proportional der elektrischen Ruh-Kraft Φ :

$$(35) \quad p = (\varepsilon - 1)\Phi.$$

ε ist die Dielektrizitätskonstante.

Hieraus wird sich nachher die Minkowskische Relation $\{C\}$, § 12, Seite 385, ergeben.

§ 9.

Die Magnetisierung bewegter Körper.

Wie der Vektor II. Art P die Polarisation, so wird Q die *Magnetisierung* bedeuten. Tatsächlich haben wir in § 6 gesehen, daß Q sich bei ruhenden Körpern auf einen Raumvektor reduziert, der als „magnetisches Moment“ zu deuten ist. Auch bei bewegten Körpern haben diese drei Komponenten von Q dieselbe Form; es sind das die Komponenten des Moments in den Koordinatenebenen yz, zx, xy ; dazu treten aber noch drei Größen, welche die Form von Momentkomponenten in den Ebenen xt, yt, zt haben. Der Vektor Q haftet also am Koordinatensystem und kann keine Eigenschaft der bewegten Materie ausdrücken. Vielmehr muß man dazu den zugehörigen Vektor I. Art

$$(36) \quad q = iwQ^* = iw \frac{\partial^2}{2} [\delta w, (\varrho_0 \dot{x})_0]^*$$

heranziehen, den wir *Ruh-Magnetisierung* nennen und der zu Q in demselben Verhältnis steht wie die *magnetische Ruh-Kraft*

$$\Psi = iw f^*$$

zu dem Feldvektor f . Die Komponenten von q sind

$$q_1 = -i \frac{\partial^2}{2} \begin{vmatrix} w_2 & w_3 & w_4 \\ \delta w_2 & \delta w_3 & \delta w_4 \\ (\varrho_0 \dot{x}_1)_0 & (\varrho_0 \dot{x}_2)_0 & (\varrho_0 \dot{x}_3)_0 \end{vmatrix}, \text{ u.s.f.}$$

Offenbar ist q normal zu w :

$$(37) \quad w\bar{q} = 0.$$

Da die zu q_1, q_2, q_3 gehörigen Determinanten je eine Vertikalreihe mit dem Index 4, d. h. mit rein imaginären Elementen haben, während die zu q_4 gehörige Determinante nur reelle Elemente hat, so sind q_1, q_2, q_3 reell und q_4 ist rein imaginär. Setzen wir also

$$q_1 = -q_x, q_2 = -q_y, q_3 = -q_z,$$

so sind das die Komponenten eines reellen Raumvektors q , und wegen (37) wird:

$$(37') \quad q_4 = -\frac{i}{c}(\mathfrak{w}_x q_x + \mathfrak{w}_y q_y + \mathfrak{w}_z q_z) = -\frac{i}{c}(\mathfrak{w} q).$$

Auch den Raumvektor q nennen wir *Ruh-Magnetisierung*; offenbar wird er für $\mathfrak{w} = 0$ mit dem in § 6 mit demselben Buchstaben bezeichneten Vektor (26) identisch.

Man kann die Gleichung (36) so deuten, daß sie die Transformation des Drehmomentes auf ein mit der Materie mitgeführtes Koordinatenkreuz ausdrückt; dabei fallen dann die Komponenten des Momentes in den Ebenen xt, yt, zt fort. q ist also ein an der Materie haftender Vektor, der ihren Zustand charakterisiert.

Der Vektor I. Art

$$\begin{aligned} w Q &= w \frac{\partial^2}{\partial^2} [\delta w, (\varrho_0 \dot{x})_0] \\ &= \frac{\partial^2}{\partial^2} \{ -(w, (\varrho_0 \dot{x})_0) \delta w + (w, \delta w) (\varrho_0 \dot{x})_0 \} \end{aligned}$$

verschwindet identisch, weil nach (5) und (7) die Vektoren $(\varrho_0 \dot{x})_0$ und δw auf w normal stehen. Wenden wir jetzt die Minkowskische Identität (§ 11, Formel (45), Seite 379)

$$[w, w Q] + [w, w Q^*]^* = (w \bar{w}) Q$$

an, so finden wir wegen (36):

$$(38) \quad [w, q]^* = -i Q.$$

Demnach drücken sich die Komponenten von Q durch die Ruh-Magnetisierung q folgendermaßen aus:

$$\begin{aligned} Q_{23} &= -\frac{q_x - \frac{\mathfrak{w}_x}{c^2}(\mathfrak{w} q)}{\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}, & Q_{31} &= -\frac{q_y - \frac{\mathfrak{w}_y}{c^2}(\mathfrak{w} q)}{\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}, & Q_{12} &= -\frac{q_z - \frac{\mathfrak{w}_z}{c^2}(\mathfrak{w} q)}{\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}, \\ Q_{14} &= -i \frac{\mathfrak{w}_y q_x - \mathfrak{w}_x q_y}{c \sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}, & Q_{24} &= -i \frac{\mathfrak{w}_z q_x - \mathfrak{w}_x q_z}{c \sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}, & Q_{34} &= -i \frac{\mathfrak{w}_x q_y - \mathfrak{w}_y q_x}{c \sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}. \end{aligned}$$

Sie setzen sich also aus den Komponenten der Raumvektoren

$$(39') \quad -\mathfrak{Q} = -\frac{q - \frac{\mathfrak{w}}{c^2}(\mathfrak{w} q)}{\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}, \quad \mathfrak{S} = \frac{[\mathfrak{w} q]}{c \sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}$$

ebenso zusammen wie die Komponenten von F aus denen von $\mathfrak{M}$ und $\mathfrak{E}$.

Der Raumvektor *Magnetisierung* $\mathfrak{Q}$ unterscheidet sich von der Ruh-Magnetisierung q nur durch Größen 2. Ordnung in $\frac{\mathfrak{w}}{c}$. Die Komponenten

von $\mathfrak{Q}$ in der Richtung $\mathfrak{w}$ und einer zu $\mathfrak{w}$ senkrechten Richtung $\mathfrak{w}$ hängen mit den entsprechenden Komponenten von q so zusammen:

$$\mathfrak{Q}_{\mathfrak{w}} = q_{\mathfrak{w}} \sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}, \quad \mathfrak{Q}_{\mathfrak{w}} = \frac{q_{\mathfrak{w}}}{\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}.$$

Der Vektor $\mathfrak{S}$ gibt, wie wir sehen werden, Anlaß zu elektrostatischen Wirkungen magnetisierter bewegter Medien; er ist das genaue Analogon zu dem „Röntgen-Vektor“.

Die für isotrope Körper gültige *Zusatzhypothese* lautet hier:

Der Raum-Zeit-Vektor I. Art Ruh-Magnetisierung q ist proportional der magnetischen Ruh-Kraft Ψ :

$$(40) \quad -q = (\mu - 1)\Psi.$$

μ ist die magnetische Permeabilität.

Hieraus wird sich nachher die Minkowskische Relation $\{D\}$, § 12, Seite 385, ergeben. Über die elektronentheoretische Deutung der Hypothesen $\{E\}$, (35), (40) gilt dasselbe, was in § 6 (S. 420) für ruhende Körper gesagt worden ist. An die Gleichung (35) hat die Theorie der Dispersion bewegter Körper anzuknüpfen.

§ 10.

Die allgemeine Beziehung zwischen den Vektoren Feldstärke und Erregung.

Die Gleichung (20), die den Raum-Zeit-Vektor f definierte, kann man jetzt nach (15') und (38) so schreiben:

$$(41) \quad \begin{aligned} f &= F + P + Q \\ &= F + [w, p] + i[w, q]^*. \end{aligned}$$

Geht man zu den reellen Raumvektoren über, so ergibt sich nach (34) und (39):

$$(42) \quad \begin{aligned} \mathfrak{m} &= \mathfrak{M} + \mathfrak{R} - \mathfrak{Q}, \\ \mathfrak{e} &= \mathfrak{E} + \mathfrak{P} + \mathfrak{S}, \end{aligned}$$

oder, ausführlich geschrieben:

$$(V') \quad \mathfrak{m} = \mathfrak{M} - \frac{q - \frac{1}{c}[w p] - \frac{w}{c^2}(w q)}{\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}},$$

$$(VI') \quad \mathfrak{e} = \mathfrak{E} + \frac{p + \frac{1}{c}[w q] - \frac{w}{c^2}(w p)}{\sqrt{1 - \frac{|\mathfrak{w}|^2}{c^2}}}.$$

Diese Formeln zusammen mit den Differentialgleichungen (I') bis (IV') und der Relation (E) stellen, genau im Sinne von H. A. Lorentz, die elektrodynamischen Vorgänge in bewegten, dielektrisch polarisierten und magnetisierten Körpern dar. Sie sind aber noch durch die Angabe des Zusammenhangs zwischen den Vektoren Ruh-Polarisation $\mathfrak{p}$ und Ruh-Magnetisierung $\mathfrak{q}$ einerseits und den wirkenden Feldstärken andererseits zu ergänzen. Ohne auf diese ergänzenden Beziehungen, die wir für isotrope Körper als „Zusatzhypothesen“ formuliert haben, einzugehen, können wir die Formeln (I') bis (VI'), (E) diskutieren und sie mit den Lorentzschen Formeln vergleichen.

Wir setzen letztere in einer zu unserem Gleichungssystem analogen Formulierung an (Enzykl. der math. Wiss., V 2, Art. 14, Seite 208, 209); die Differentialgleichungen lauten:

$$(III'') \text{ u. (XXVIII) } \operatorname{curl} \mathfrak{H} - \frac{1}{c} \frac{\partial \mathfrak{D}}{\partial t} = \frac{1}{c} (\mathfrak{s} + \operatorname{curl} [\mathfrak{P} \mathfrak{w}]),$$

$$(I'') \quad \operatorname{div} \mathfrak{D} = \mathfrak{q},$$

$$(IV'') \quad \operatorname{curl} \mathfrak{E} + \frac{1}{c} \frac{\partial \mathfrak{B}}{\partial t} = 0,$$

$$(V'') \quad \operatorname{div} \mathfrak{B} = 0.$$

Dazu kommen die unseren Gleichungen (V'), (VI') entsprechenden Relationen:

$$(XXX) \quad \mathfrak{H} = \mathfrak{B} - \mathfrak{Q},$$

$$(XXV) \text{ u. (XXI) } \mathfrak{D} = \mathfrak{E} + \mathfrak{P}.$$

Dabei haben wir im allgemeinen die Lorentzschen Bezeichnungen beibehalten; nur ist die Magnetisierung mit $\mathfrak{Q}$ (statt mit $\mathfrak{M}$) bezeichnet und es sind Leitungsstrom $\mathfrak{J}$ und Konvektionsstrom $\mathfrak{K}$ zu dem Stromvektor $\mathfrak{s}$ zusammengefaßt.

Offenbar gehen die Lorentzschen Differentialgleichungen in die unseren über, wenn wir die Lorentzschen Vektoren $\mathfrak{E}$, $\mathfrak{H} - \frac{1}{c} [\mathfrak{P} \mathfrak{w}]$, $\mathfrak{D}$, $\mathfrak{B}$ bzw. mit $\mathfrak{E}$, $\mathfrak{m}$, $\mathfrak{e}$, $\mathfrak{M}$ identifizieren). Die beiden weiteren Gleichungen lassen sich für beliebige Körper nicht zur Übereinstimmung bringen.*

Bei nichtmagnetisierten Körpern ($\mathfrak{q} = 0$, $\mathfrak{Q} = 0$) läßt sich aber Übereinstimmung erreichen, wenn man in (V'), (VI') alle in $\frac{w}{c}$ quadratischen Glieder vernachlässigt; dann bleibt nämlich (wegen (34') und (39')):

*) Minkowski identifiziert in § 9 der „Grundgleichungen“ $\mathfrak{H}$ mit $\mathfrak{m}$; daraus entspringt es, daß er die Übereinstimmung der Lorentzschen Formeln für nichtmagnetisierte Körper mit dem Relativitätsprinzip dem Umstande zuschreiben muß, daß Lorentz die Bedingung des Nichtmagnetisiertseins in einer dem Prinzipie widersprechenden Weise ansetzt. Unsere Betrachtungsweise führt von selbst auf obige Zuordnung der Vektoren, wobei die schließliche Übereinstimmung nicht auf die Kompensation zweier Widersprüche zurückgeführt zu werden braucht.

$$(42') \quad \begin{aligned} \mathfrak{m} &= \mathfrak{M} - \mathfrak{Q} + \frac{1}{c} [\mathfrak{w} \mathfrak{P}], \\ \mathfrak{e} &= \mathfrak{E} + \mathfrak{P} + \frac{1}{c} [\mathfrak{w} \mathfrak{Q}], \end{aligned}$$

und die Formeln gehen für $\mathfrak{Q} = 0$ bei der oben angegebenen Zuordnung der Lorentzschen Vektoren zu den unseren in die Lorentzschen Relationen über.

Hieraus erhellt, daß unsere Bezeichnung des Vektors $\mathfrak{R}$, der bis auf Glieder 2. Ordnung durch $\frac{1}{c} [\mathfrak{w} \mathfrak{P}]$ gegeben ist, als „Röntgen-Vektor“ gerechtfertigt ist. Denn er gibt wie bei Lorentz zur Entstehung des Röntgenstroms Anlaß, und zwar in einer mit Eichenwalds Versuchen übereinstimmenden Weise (Enzykl., V 2, Art. 14, S. 210). Die Formeln (V'), (VI') sowohl, als auch die Näherungsformeln (42') unterscheiden sich von den Lorentzschen durch die volle Symmetrie*) zwischen den elektrischen und den magnetischen Größen. Sie beruht vor allem auf dem Vorhandensein des Vektors $\mathfrak{S}$, der bis auf Glieder 2. Ordnung durch $\frac{1}{c} [\mathfrak{w} \mathfrak{Q}]$ gegeben ist; dieser entspricht genau dem Röntgen-Vektor und zeigt elektrostatische Wirkungen bewegter, magnetisierter Körper an. *Erscheinungen, die durch diesen Vektor $\mathfrak{S}$ zu erklären sind, sind meines Wissens bislang experimentell noch nicht festgestellt worden*; es lassen sich aber unschwer Versuchsanordnungen angeben, die eine experimentelle Untersuchung über das Vorhandensein dieser Wirkung, die in $\frac{m}{c}$ von erster Ordnung ist, gestatten würden. Dieselbe ist nicht zu verwechseln mit der Erscheinung, daß ein Leiter, der sich im magnetischen Felde bewegt, in ruhenden Strombahnen Leitungsströme hervorruft (z. B. im Falle der sog. unipolaren Induktion); diese letztere Erscheinung ist eine Konsequenz der Formeln (E), S. 422, und man sieht leicht, daß sie bis auf Glieder 2. Ordnung ebenso wie in den Theorien von Hertz und Lorentz durch den Vektor $\frac{1}{c} [\mathfrak{w} \mathfrak{M}]$ bestimmt ist**), dessen Betrag leicht wesentlich größer gemacht werden kann als der des Vektors $\frac{1}{c} [\mathfrak{w} \mathfrak{Q}]$, weil $\mathfrak{Q}$ bei den meisten Körpern nur wenig von Null verschieden ist.

*) Die Hertzschen Grundgleichungen der Elektrodynamik bewegter Körper zeigen eine analoge Symmetrie, widersprechen aber bekanntlich dem Experiment und sind mit dem Relativitätsprinzip unvereinbar.

Ein Vergleich unserer Formeln mit den Grundgleichungen von E. Cohn erübrigt sich, da demselben der § 10 der Minkowskischen „Grundgleichungen“ gewidmet ist.

**) Vgl. Lorentz, Enzyklopädie, V 1, Art. 13, S. 100 u. S. 238, 239. M. Abraham, Theorie der Elektrizität, Bd. I, § 86, 87, S. 398—409, und Bd. II (Elektromagnetische Theorie der Strahlung), 2. Aufl. 1908, § 36, S. 300—302.

Was die Glieder 2. Ordnung in $\frac{w}{c}$ anbetrifft, so scheint wenig Aussicht vorhanden zu sein, sie bei irdischen Experimenten aufzufinden.

Wir wollen nun noch zeigen, daß bei isotropen, nichtdispargierenden Körpern unsere Zusatzhypothesen (35) und (40) auf die Minkowskischen Formeln $\{C\}$, $\{D\}$ zurückführen.

Nach § 8, (33) ist $wP = -p$ und nach § 9, S. 425, ist identisch $wQ = 0$. Daher folgt aus (20)

$$wf = wF - p.$$

Da nun $\Phi = -wF$ ist, so ergibt sich aus (35):

$$\{C\} \quad wf = \varepsilon wF$$

oder

$$(C) \quad e + \frac{1}{c}[wm] = \varepsilon \left(\mathfrak{E} + \frac{1}{c}[wM] \right).$$

Für die zu f, F, P, Q dualen Vektoren f^*, F^*, P^*, Q^* gilt nach (20) die Gleichung:

$$f^* = F^* + P^* + Q^*.$$

Nun ist offenbar

$$wP^* = w[w, p]^* = 0,$$

und nach § 9, (36) ist

$$wQ^* = -iq.$$

Daher wird

$$wf^* = wF^* - iq.$$

Da nun $\Psi = iw f^*$ ist, so folgt aus (40):

$$\{D\} \quad wF^* = \mu w f^*$$

oder

$$(D) \quad M - \frac{1}{c}[w\mathfrak{E}] = \mu \left(m - \frac{1}{c}[we] \right).$$

Damit sind wir zu dem vollständigen Formelsystem Minkowskis gelangt.

Wie schon hervorgehoben, gibt die größere Allgemeinheit unserer Formeln (42) die Möglichkeit an die Hand, komplizierteren Eigenschaften der Substanzen durch die Theorie gerecht zu werden; dazu sind nur die „Zusatzhypothesen“ (35), (40) durch andere zu ersetzen. Die Dispersion in bewegten Medien wird man z. B. erhalten, wenn man die Gleichung (35) durch geeignete Glieder so erweitert, daß der Vektor p Schwingungen träger Massen darzustellen fähig ist. Vorstehende Theorie, die eine Mittelstellung einnimmt zwischen der ursprünglichen rein phänomenologischen Methode Minkowskis und dem bis ins Detail der Elektronenbewegung eindringenden Verfahren von Lorentz, scheint also einerseits schmiegsam genug zu sein, die mannigfaltigen Eigenarten der Substanzen zum Ausdruck zu bringen, andererseits besitzt sie ein anschauliches Fundament, das ihrer bequemen Handhabung förderlich ist.

Inhaltsübersicht.

	Seite
Einleitung	406
§ 1. Bezeichnungen	409
§ 2. Die Zerlegung der elektrischen Strömung	410
§ 3. Die Darstellung der variirten Strömung	411
§ 4. Die Reihenentwicklung der variirten Strömung	415
§ 5. Formale Herstellung der in der bewegten Materie gültigen Differentialgleichungen	416
§ 6. Ruhende Körper	418
§ 7. Der Leitungsstrom in bewegten Körpern	420
§ 8. Die dielektrische Polarisirung in bewegten Körpern	422
§ 9. Die Magnetisirung bewegter Körper	424
§ 10. Die allgemeine Beziehung zwischen den Vektoren Feldstärke und Erregung	426

XXXII.

Raum und Zeit.

(Vortrag, gehalten auf der 80. Naturforscher-Versammlung zu Köln am 21. Sept. 1908.)
(Physikalische Zeitschrift, 10. Jahrgang, 1909, S. 104—111 und Jahresbericht der Deutschen Mathematiker-Vereinigung, Bd. 18, S. 75—88; auch als Sonderabdruck erschienen, Leipzig, B. G. Teubner 1909.)*)

M. H.! Die Anschauungen über Raum und Zeit, die ich Ihnen entwickeln möchte, sind auf experimentell-physikalischem Boden erwachsen. Darin liegt ihre Stärke. Ihre Tendenz ist eine radikale. Von Stund an sollen Raum für sich und Zeit für sich völlig zu Schatten herabsinken, und nur noch eine Art Union der beiden soll Selbständigkeit bewahren.

I.

Ich möchte zunächst ausführen, wie man von der gegenwärtig angenommenen Mechanik wohl durch eine rein mathematische Überlegung zu veränderten Ideen über Raum und Zeit kommen könnte. Die Gleichungen der Newtonschen Mechanik zeigen eine zweifache Invarianz. Einmal bleibt ihre Form erhalten, wenn man das zugrunde gelegte räumliche Koordinatensystem einer beliebigen *Lagenveränderung* unterwirft, zweitens, wenn man es in seinem Bewegungszustande verändert, nämlich ihm irgendeine *gleichförmige Translation* aufprägt; auch spielt der Nullpunkt der Zeit keine Rolle. Man ist gewohnt, die Axiome der Geometrie als erledigt anzusehen, wenn man sich reif für die Axiome der Mechanik fühlt, und deshalb werden jene zwei Invarianzen wohl selten in einem Atemzuge genannt. Jede von ihnen bedeutet eine gewisse Gruppe von Transformationen in sich für die Differentialgleichungen der Mechanik. Die Existenz der ersteren Gruppe sieht man als einen fundamentalen Charakter des Raumes an. Die zweite Gruppe straft man am liebsten mit Verachtung, um leichten Sinnes darüber hinwegzukommen, daß man von den physikalischen Erscheinungen her niemals entscheiden kann, ob der als ruhend vorausgesetzte Raum sich nicht am Ende in einer gleichförmigen Translation befindet. So führen jene zwei Gruppen ein völlig getrenntes Dasein nebeneinander. Ihr gänzlich heterogener Charakter mag davon abgeschreckt haben, sie zu komponieren. Aber gerade die komponierte volle Gruppe als Ganzes gibt uns zu denken auf.

*) Dem Vortrag war im Sonderabdruck ein Bildnis Minkowskis als Titelbild beigegeben. (Anm. d. Herausg.)

Wir wollen uns die Verhältnisse graphisch zu veranschaulichen suchen. Es seien x, y, z rechtwinklige Koordinaten für den Raum und t bezeichne die Zeit. Gegenstand unserer Wahrnehmung sind immer nur Orte und Zeiten verbunden. Es hat niemand einen Ort anders bemerkt als zu einer Zeit, eine Zeit anders als an einem Orte. Ich respektiere aber noch das Dogma, daß Raum und Zeit je eine unabhängige Bedeutung haben. Ich will einen Raumpunkt zu einem Zeitpunkt, d. i. ein Wertsystem x, y, z, t einen *Weltpunkt* nennen. Die Mannigfaltigkeit aller denkbaren Wertsysteme x, y, z, t soll die *Welt* heißen. Ich könnte mit kühner Kreide vier Weltachsen auf die Tafel werfen. Schon *eine* gezeichnete Achse besteht aus lauter schwingenden Molekülen und macht zudem die Reise der Erde im All mit, gibt also bereits genug zu abstrahieren auf; die mit der Anzahl 4 verbundene etwas größere Abstraktion tut dem Mathematiker nicht wehe. Um nirgends eine gähnende Leere zu lassen, wollen wir uns vorstellen, daß allerorten und zu jeder Zeit etwas Wahrnehmbares vorhanden ist. Um nicht Materie oder Elektrizität zu sagen, will ich für dieses Etwas das Wort Substanz brauchen. Wir richten unsere Aufmerksamkeit auf den im Weltpunkt x, y, z, t vorhandenen substantiellen Punkt und stellen uns vor, wir sind imstande, diesen substantiellen Punkt zu jeder anderen Zeit wiederzuerkennen. Einem Zeitelement dt mögen die Änderungen dx, dy, dz der Raumkoordinaten dieses substantiellen Punktes entsprechen. Wir erhalten alsdann als Bild sozusagen für den ewigen Lebenslauf des substantiellen Punktes eine Kurve in der Welt, eine *Weltlinie*, deren Punkte sich eindeutig auf den Parameter t von $-\infty$ bis $+\infty$ beziehen lassen. Die ganze Welt erscheint aufgelöst in solche Weltlinien, und ich möchte sogleich vorwegnehmen, daß meiner Meinung nach die physikalischen Gesetze ihren vollkommensten Ausdruck als Wechselbeziehungen unter diesen Weltlinien finden dürften.

Durch die Begriffe Raum und Zeit fallen die x, y, z -Mannigfaltigkeit $t=0$ und ihre zwei Seiten $t>0$ und $t<0$ auseinander. Halten wir der Einfachheit wegen den Nullpunkt von Raum und Zeit fest, so bedeutet die zuerst genannte Gruppe der Mechanik, daß wir die x, y, z -Achsen in $t=0$ einer beliebigen Drehung um den Nullpunkt unterwerfen dürfen, entsprechend den homogenen linearen Transformationen des Ausdrucks

$$x^2 + y^2 + z^2$$

in sich. Die zweite Gruppe aber bedeutet, daß wir, ebenfalls ohne den Ausdruck der mechanischen Gesetze zu verändern,

$$x, y, z, t \quad \text{durch} \quad x - \alpha t, y - \beta t, z - \gamma t, t$$

mit irgendwelchen Konstanten α, β, γ ersetzen dürfen. Der Zeitachse kann

hiernach eine völlig beliebige Richtung nach der oberen halben Welt $t > 0$ gegeben werden. Was hat nun die Forderung der Orthogonalität im Raume mit dieser völligen Freiheit der Zeitachse nach oben hin zu tun?

Die Verbindung herzustellen, nehmen wir einen positiven Parameter c und betrachten das Gebilde

$$c^2 t^2 - x^2 - y^2 - z^2 = 1.$$

Es besteht aus zwei durch $t = 0$ getrennten Schalen nach Analogie eines zweischaligen Hyperboloids. Wir betrachten die Schale im Gebiete $t > 0$ und wir fassen jetzt diejenigen homogenen linearen Transformationen von x, y, z, t in vier neue Variable x', y', z', t' auf, wobei der Ausdruck dieser Schale in den neuen Variablen entsprechend wird. Zu diesen Transformationen gehören offenbar die Drehungen des Raumes um den Nullpunkt. Ein volles Verständnis der übrigen jener Transformationen erhalten wir hernach bereits, wenn wir eine solche unter ihnen ins Auge fassen, bei der y und z ungeändert

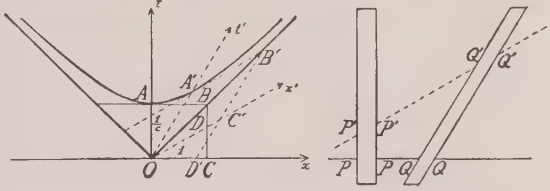


Fig. 1.

bleiben. Wir zeichnen (Fig. 1) den Durchschnitt jener Schale mit der Ebene der x - und der t -Achse, den oberen Ast der Hyperbel $c^2 t^2 - x^2 = 1$, mit seinen Asymptoten. Ferner werde ein beliebiger Radiusvektor OA' dieses Hyperbelastes vom Nullpunkte O aus eingetragen, die Tangente in A' an die Hyperbel bis zum Schnitte B' mit der Asymptote rechts gelegt, $OA'B'$ zum Parallelogramm $OA'B'C'$ vervollständigt, endlich für das spätere noch $B'C'$ bis zum Schnitt D' mit der x -Achse durchgeführt. Nehmen wir nun OC' und OA' als Achsen für Parallelkoordinaten x', t' mit den Maßstäben $OC' = 1$, $OA' = 1/c$, so erlangt jener Hyperbelast wieder den Ausdruck $c^2 t'^2 - x'^2 = 1$, $t' > 0$, und der Übergang von x, y, z, t zu x', y, z, t' ist eine der fraglichen Transformationen. Wir nehmen nun zu den charakterisierten Transformationen noch die beliebigen Verschiebungen des Raum- und Zeit-Nullpunktes hinzu und konstituieren damit eine offenbar noch von dem Parameter c abhängige Gruppe von Transformationen, die ich mit G_c bezeichne.

Lassen wir jetzt c ins Unendliche wachsen, also $1/c$ nach Null konvergieren, so leuchtet an der beschriebenen Figur ein, daß der Hyperbelast sich immer mehr der x -Achse anschmiegt, der Asymptotenwinkel sich zu einem gestreckten verbreitert, jene spezielle Transformation in der Grenze sich in eine solche verwandelt, wobei die t' -Achse eine beliebige Richtung nach oben haben kann und x' immer genauer

sich an x annähert. Mit Rücksicht hierauf ist klar, daß aus der Gruppe G_c in der Grenze für $c = \infty$, also als Gruppe G_∞ , eben jene zu der Newtonschen Mechanik gehörige volle Gruppe wird. Bei dieser Sachlage, und da G_c mathematisch verständlicher ist als G_∞ , hätte wohl ein Mathematiker in freier Phantasie auf den Gedanken verfallen können, daß am Ende die Naturerscheinungen tatsächlich eine Invarianz nicht bei der Gruppe G_∞ , sondern vielmehr bei einer Gruppe G_c mit bestimmtem endlichen, nur in den gewöhnlichen Maßeinheiten *äußerst großen* c besitzen. Eine solche Ahnung wäre ein außerordentlicher Triumph der reinen Mathematik gewesen. Nun, da die Mathematik hier nur mehr Treppenwitz bekundet, bleibt ihr doch die Genugtuung, daß sie dank ihren glücklichen Antezedenzen mit ihren in freier Fernsicht geschärften Sinnen die tiefgreifenden Konsequenzen einer solchen Ummodelung unserer Naturauffassung auf der Stelle zu erfassen vermag.

Ich will sogleich bemerken, um welchen Wert für c es sich schließlich handeln wird. Für c wird die *Fortpflanzungsgeschwindigkeit des Lichtes im leeren Raume* eintreten. Um weder vom Raum, noch von Leere zu sprechen, können wir diese Größe wieder als das Verhältnis der elektromagnetischen und der elektrostatischen Einheit der Elektrizitätsmenge kennzeichnen.

Das Bestehen der Invarianz der Naturgesetze für die bezügliche Gruppe G_c würde nun so zu fassen sein:

Man kann aus der Gesamtheit der Naturerscheinungen durch sukzessiv gesteigerte Approximationen immer genauer ein Bezugssystem x, y, z und t , Raum und Zeit, ableiten, mittels dessen diese Erscheinungen sich dann nach bestimmten Gesetzen darstellen. Dieses Bezugssystem ist dabei aber durch die Erscheinungen keineswegs eindeutig festgelegt. *Man kann das Bezugssystem noch entsprechend den Transformationen der genannten Gruppe G_c beliebig verändern, ohne daß der Ausdruck der Naturgesetze sich dabei verändert.*

Z. B. kann man der beschriebenen Figur entsprechend auch t' Zeit benennen, muß dann aber im Zusammenhange damit notwendig den Raum durch die Mannigfaltigkeit der drei Parameter x', y, z definieren, wobei nun die physikalischen Gesetze mittels x', y, z, t' sich genau ebenso ausdrücken würden, wie mittels x, y, z, t . Hiernach würden wir dann in der Welt nicht mehr *den* Raum, sondern unendlich viele Räume haben, analog wie es im dreidimensionalen Raume unendlich viele Ebenen gibt. Die dreidimensionale Geometrie wird ein Kapitel der vierdimensionalen Physik. Sie erkennen, weshalb ich am Eingange sagte, Raum und Zeit sollen zu Schatten herabsinken und nur eine Welt an sich bestehen.

II.

Nun ist die Frage, welche Umstände zwingen uns die veränderte Auffassung von Raum und Zeit auf, widerspricht sie tatsächlich niemals den Erscheinungen, endlich gewährt sie Vorteile für die Beschreibung der Erscheinungen?

Bevor wir hierauf eingehen, sei eine wichtige Bemerkung vorangestellt. Haben wir Raum und Zeit irgendwie individualisiert, so entspricht einem ruhenden substantiellen Punkte als Weltlinie eine zur t -Achse parallele Gerade, einem gleichförmig bewegten substantiellen Punkte eine gegen die t -Achse geneigte Gerade, einem ungleichförmig bewegten substantiellen Punkte eine irgendwie gekrümmte Weltlinie. Fassen wir in einem beliebigen Weltpunkte x, y, z, t die dort durchlaufende Weltlinie auf und finden wir sie dort parallel mit irgendeinem Radiusvektor OA' der vorhin genannten hyperboloidischen Schale, so können wir OA' als neue Zeitachse einführen, und bei den damit gegebenen neuen Begriffen von Raum und Zeit erscheint die Substanz in dem betreffenden Weltpunkte als ruhend. Wir wollen nun dieses fundamentale Axiom einführen:

Die in einem beliebigen Weltpunkte vorhandene Substanz kann stets bei geeigneter Festsetzung von Raum und Zeit als ruhend aufgefaßt werden.

Das Axiom bedeutet, daß in jedem Weltpunkte stets der Ausdruck

$$c^2 dt^2 - dx^2 - dy^2 - dz^2$$

positiv ausfällt oder, was damit gleichbedeutend ist, daß jede Geschwindigkeit v stets kleiner als c ausfällt. Es würde danach für alle substantiellen Geschwindigkeiten c als obere Grenze bestehen und hierin eben die tiefere Bedeutung der Größe c liegen. In dieser anderen Fassung hat das Axiom beim ersten Eindruck etwas Mißfälliges. Es ist aber zu bedenken, daß nun eine modifizierte Mechanik Platz greifen wird, in der die Quadratwurzel aus jener Differentialverbindung zweiten Grades eingeht, so daß Fälle mit Überlichtgeschwindigkeit nur mehr eine Rolle spielen werden, etwa wie in der Geometrie Figuren mit imaginären Koordinaten.

Der Anstoß und wahre Beweggrund für die Annahme der Gruppe G_0 nun kam daher, daß die Differentialgleichung für die Fortpflanzung von Lichtwellen im leeren Raume jene Gruppe G_0 besitzt*). Andererseits hat der Begriff starrer Körper nur in einer Mechanik mit der Gruppe G_∞ einen Sinn. Hat man nun eine Optik mit G_0 , und gäbe es andererseits

*) Eine wesentliche Anwendung dieser Tatsache findet sich bereits bei W. Voigt, Nachrichten der K. Gesellschaft der Wissenschaften zu Göttingen, mathematisch-physikalische Klasse, 1887, S. 41.

starre Körper, so ist leicht abzusehen, daß durch die zwei zu G_c und zu G_∞ gehörigen hyperboloidischen Schalen *eine* t -Richtung ausgezeichnet sein würde, und das würde weiter die Konsequenz haben, daß man an geeigneten starren optischen Instrumenten im Laboratorium einen Wechsel der Erscheinungen bei verschiedener Orientierung gegen die Fortschreitungsrichtung der Erde müßte wahrnehmen können. Alle auf dieses Ziel gerichteten Bemühungen, insbesondere ein berühmter Interferenzversuch von Michelson, hatten jedoch ein negatives Ergebnis. Um eine Erklärung hierfür zu gewinnen, bildete H. A. Lorentz eine Hypothese, deren Erfolg eben in der Invarianz der Optik für die Gruppe G_c liegt. Nach Lorentz soll jeder Körper, der eine Bewegung besitzt, in Richtung der Bewegung eine Verkürzung erfahren haben und zwar bei einer Geschwindigkeit v im Verhältnisse

$$1 : \sqrt{1 - \frac{v^2}{c^2}}.$$

Diese Hypothese klingt äußerst phantastisch. Denn die Kontraktion ist nicht etwa als Folge von Widerständen im Äther zu denken, sondern rein als Geschenk von oben, als Begleitumstand des Umstandes der Bewegung.

Ich will nun an unserer Figur zeigen, daß die Lorentzsche Hypothese völlig äquivalent ist mit der neuen Auffassung von Raum und Zeit, wodurch sie viel verständlicher wird. Abstrahieren wir der Einfachheit wegen von y und z und denken uns eine räumlich eindimensionale Welt, so sind ein wie die t -Achse aufrechter und ein gegen die t -Achse geneigter Parallelstreifen (siehe Fig. 1) Bilder für den Verlauf eines ruhenden, bezüglich eines gleichförmig bewegten Körpers, der jedesmal eine konstante räumliche Ausdehnung behält. Ist OA' parallel dem zweiten Streifen, so können wir t' als Zeit und x' als Raumkoordinate einführen, und es erscheint dann der zweite Körper als ruhend, der erste als gleichförmig bewegt. Wir nehmen nun an, daß der erste Körper als ruhend aufgefaßt die Länge l hat, d. h. der Querschnitt PP des ersten Streifens auf der x -Achse $= l \cdot OC$ ist, wo OC den Einheitsmaßstab auf der x -Achse bedeutet, und daß andererseits der zweite Körper *als ruhend aufgefaßt* die gleiche Länge l hat; letzteres heißt dann, daß der *parallel der x' -Achse* gemessene Querschnitt des zweiten Streifens, $Q'Q' = l \cdot OC'$ ist. Wir haben nunmehr in diesen zwei Körpern Bilder von zwei *gleichen* Lorentzschen Elektronen, einem ruhenden und einem gleichförmig bewegten. Halten wir aber an den ursprünglichen Koordinaten x, t fest, so ist als Ausdehnung des zweiten Elektrons der Querschnitt QQ seines zugehörigen Streifens *parallel der x -Achse* anzugeben. Nun ist offenbar, da $Q'Q' = l \cdot OC'$

ist, $QQ = l \cdot OD'$. Eine leichte Rechnung ergibt, wenn dx/dt für den zweiten Streifen $= v$ ist, $OD' = OC \cdot \sqrt{1 - \frac{v^2}{c^2}}$, also auch $PP : QQ = 1 : \sqrt{1 - \frac{v^2}{c^2}}$. Dies ist aber der Sinn der Lorentz'schen Hypothese von der Kontraktion der Elektronen bei Bewegung. Fassen wir andererseits das zweite Elektron als ruhend auf, adoptieren also das Bezugssystem x', t' , so ist als Länge des ersten der Querschnitt $P'P'$ seines Streifens parallel OC' zu bezeichnen, und wir würden in genau dem nämlichen Verhältnisse das erste Elektron gegen das zweite verkürzt finden; denn es ist in der Figur

$$PP' : Q'Q' = OD : OC' = OD' : OC = QQ : PP.$$

Lorentz nannte die Verbindung t' von x und t *Ortszeit* des gleichförmig bewegten Elektrons und verwandte eine physikalische Konstruktion dieses Begriffs zum besseren Verständnis der Kontraktionshypothese. Jedoch scharf erkannt zu haben, daß die Zeit des einen Elektrons ebenso gut wie die des anderen ist, d. h. daß t und t' gleich zu behandeln sind, ist erst das Verdienst von A. Einstein*). Damit war nun zunächst die Zeit als ein durch die Erscheinungen eindeutig festgelegter Begriff abgesetzt. An dem Begriffe des Raumes rüttelten weder Einstein noch Lorentz, vielleicht deshalb nicht, weil bei der genannten speziellen Transformation, wo die x', t' -Ebene sich mit der x, t -Ebene deckt, eine Deutung möglich ist, als sei die x -Achse des Raumes in ihrer Lage erhalten geblieben. Über den Begriff des Raumes in entsprechender Weise hinwegzuschreiten, ist auch wohl nur als Verwegenheit mathematischer Kultur einzutaxieren. Nach diesem zum wahren Verständnis der Gruppe G_c jedoch unerläßlichen weiteren Schritt aber scheint mir das Wort *Relativitätspostulat* für die Forderung einer Invarianz bei der Gruppe G_c sehr matt. Indem der Sinn des Postulats wird, daß durch die Erscheinungen nur die in Raum und Zeit vierdimensionale Welt gegeben ist, aber die Projektion in Raum und in Zeit noch mit einer gewissen Freiheit vorgenommen werden kann, möchte ich dieser Behauptung eher den Namen *Postulat der absoluten Welt* (oder kurz *Weltpostulat*) geben.

III.

Durch das Weltpostulat wird eine gleichartige Behandlung der vier Bestimmungsstücke x, y, z, t möglich. Dadurch gewinnen, wie ich jetzt

*) A. Einstein, Annalen der Physik, Bd. 17, 1905, S. 891; Jahrbuch der Radioaktivität und Elektronik, Bd. 4, 1907, S. 411.

ausführen will, die Formen, unter denen die physikalischen Gesetze sich abspielen, an Verständlichkeit. Vor allem erlangt der Begriff der *Beschleunigung* ein scharf hervortretendes Gepräge.

Ich werde mich einer geometrischen Ausdrucksweise bedienen, die sich sofort darbietet, indem man im Tripel x, y, z stillschwei-

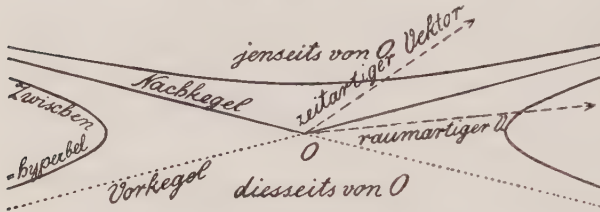


Fig. 2.

gend von z abstrahiert. Einen beliebigen Weltpunkt O denke ich zum Raum-Zeit-Nullpunkt gemacht. Der Kegel

$$c^2 t^2 - x^2 - y^2 - z^2 = 0$$

mit O als Spitze (Fig. 2) besteht aus zwei Teilen, einem mit Werten $t < 0$, einem anderen mit Werten $t > 0$. Der erste, der *Vorkegel* von O , besteht, sagen wir, aus allen Weltpunkten, die „Licht nach O senden“, der zweite, der *Nachkegel* von O , aus allen Weltpunkten, die „Licht von O empfangen.“ Das vom Vorkegel allein begrenzte Gebiet mag *diesseits von O* , das vom Nachkegel allein begrenzte *jenseits von O* heißen. Jenseits O fällt die schon betrachtete hyperboloidische Schale

$$F = c^2 t^2 - x^2 - y^2 - z^2 = 1, \quad t > 0.$$

Das Gebiet *zwischen den Kegeln* wird erfüllt von den einschaligen hyperboloidischen Gebilden

$$-F = x^2 + y^2 + z^2 - c^2 t^2 = k^2$$

zu allen konstanten positiven Werten k^2 . Wichtig sind für uns die Hyperbeln mit O als Mittelpunkt, die auf den letzteren Gebilden liegen. Die einzelnen Äste dieser Hyperbeln mögen kurz die *Zwischenhyperbeln zum Zentrum O* heißen. Ein solcher Hyperbelast würde, als Weltlinie eines substantiellen Punktes gedacht, eine Bewegung repräsentieren, die für $t = -\infty$ und $t = +\infty$ asymptotisch auf die Lichtgeschwindigkeit c ansteigt.

Nennen wir in Analogie zum Vektorbegriff im Raume jetzt eine gerichtete Strecke in der Mannigfaltigkeit der x, y, z, t einen *Vektor*, so haben wir zu unterscheiden zwischen den *zeitartigen* Vektoren mit Richtungen von O nach der Schale $+F=1$, $t > 0$ und den *raumartigen* Vektoren mit Richtungen von O nach $-F=1$. Die Zeitachse kann jedem Vektor der ersteren Art parallel laufen. Ein jeder Weltpunkt zwischen Vorkegel und Nachkegel von O kann durch das Bezugssystem als *gleichzeitig* mit O , aber ebensogut auch als *früher* als O oder als *später* als O eingerichtet werden. Jeder Weltpunkt diesseits O ist not-

wendig stets früher, jeder Weltpunkt jenseits O notwendig stets später als O . Dem Grenzübergang zu $c = \infty$ würde ein völliges Zusammenklappen des keilförmigen Einschnittes zwischen den Kegeln in die ebene Mannigfaltigkeit $t = 0$ entsprechen. In den gezeichneten Figuren ist dieser Einschnitt absichtlich mit verschiedener Breite angelegt.

Einen beliebigen Vektor, wie von O nach x, y, z, t , zerlegen wir in die vier *Komponenten* x, y, z, t . Sind die Richtungen zweier Vektoren beziehungsweise die eines Radiusvektors OR von O an eine der Flächen $\mp F = 1$ und dazu einer Tangente RS im Punkte R der betreffenden Fläche, so sollen die Vektoren *normal* zueinander heißen. Danach ist

$$c^2 tt_1 - xx_1 - yy_1 - zz_1 = 0$$

die Bedingung dafür, daß die Vektoren mit den Komponenten x, y, z, t und x_1, y_1, z_1, t_1 *normal* zueinander sind.

Für die *Beträge* von Vektoren der verschiedenen Richtungen sollen die *Einheitsmaßstäbe* dadurch fixiert sein, daß einem raumartigen Vektor von O nach $-F = 1$ stets der Betrag 1, einem zeitartigen Vektor von O nach $+F = 1$, $t > 0$ stets der Betrag $1/c$ zugeschrieben wird.

Denken wir uns nun in einem Weltpunkte $P(x, y, z, t)$ die dort durchlaufende Weltlinie eines substantiellen Punktes, so entspricht danach dem zeitartigen Vektorelement dx, dy, dz, dt im Fortgang der Linie der Betrag

$$d\tau = \frac{1}{c} \sqrt{c^2 dt^2 - dx^2 - dy^2 - dz^2}.$$

Das Integral $\int d\tau = \tau$ dieses Betrages auf der Weltlinie von irgendeinem fixierten Ausgangspunkte P_0 bis zu dem variablen Endpunkte P geführt, nennen wir die *Eigenzeit* des substantiellen Punktes in P . Auf der Weltlinie betrachten wir x, y, z, t , d. s. die Komponenten des Vektors OP , als Funktionen der Eigenzeit τ , bezeichnen deren erste Differentialquotienten nach τ mit $\dot{x}, \dot{y}, \dot{z}, \dot{t}$, deren zweite Differentialquotienten nach τ mit $\ddot{x}, \ddot{y}, \ddot{z}, \ddot{t}$ und nennen die zugehörigen Vektoren, die Ableitung des Vektors OP nach τ den *Bewegungsvektor* in P und die Ableitung dieses Bewegungsvektors nach τ den *Beschleunigungsvektor* in P . Dabei gilt

$$c^2 \dot{t}^2 - \dot{x}^2 - \dot{y}^2 - \dot{z}^2 = c^2,$$

$$c^2 \ddot{t} - \ddot{x} \dot{x} - \ddot{y} \dot{y} - \ddot{z} \dot{z} = 0,$$

d. h. der Bewegungsvektor ist der zeitartige Vektor in Richtung der Weltlinie in P vom Betrage 1 und der Beschleunigungsvektor in P ist normal zum Bewegungsvektor in P , also jedenfalls ein raumartiger Vektor.

Nun gibt es, wie man leicht einsieht, einen bestimmten Hyperbelast, der mit der Weltlinie in P drei unendlich benachbarte Punkte

gemein hat und dessen Asymptoten Erzeugende eines Vorkegels und eines Nachkegels sind (siehe unten Fig. 3). Dieser Hyperbelast heie die *Krmmungshyperbel* in P . Ist M das Zentrum dieser Hyperbel, so handelt es sich also hier um eine Zwischenhyperbel zum Zentrum M . Es sei ρ der Betrag des Vektors MP , so erkennen wir den Beschleunigungsvektor in P als den Vektor in Richtung MP vom Betrage c^2/ρ .

Sind $\ddot{x}, \ddot{y}, \ddot{z}, \ddot{t}$ smtlich Null, so reduziert sich die Krmmungshyperbel auf die in P die Weltlinie berhrende Gerade, und es ist $\rho = \infty$ zu setzen.

IV.

Um darzutun, da die Annahme der Gruppe G_c fr die physikalischen Gesetze nirgends zu einem Widerspruche fhrt, ist es unumgnglich, eine Revision der gesamten Physik auf Grund der Voraussetzung dieser Gruppe vorzunehmen. Diese Revision ist bereits in einem gewissen Umfange erfolgreich geleistet fr Fragen der Thermodynamik und Wrmestrahlung*), fr die elektromagnetischen Vorgnge, endlich fr die Mechanik unter Aufrechterhaltung des Massenbegriffes**).

Fr letzteres Gebiet ist vor allem die Frage aufzuwerfen: Wenn eine Kraft mit den Komponenten X, Y, Z nach den Raumachsen in einem Weltpunkte $P(x, y, z, t)$ angreift, wo der Bewegungsvektor $\dot{x}, \dot{y}, \dot{z}, \dot{t}$ ist, als welche Kraft ist diese Kraft bei einer beliebigen nderung des Bezugssystemes aufzufassen? Nun existieren gewisse erprobte Anstze ber die ponderomotorische Kraft im elektromagnetischen Felde in den Fllen, wo die Gruppe G_c unzweifelhaft zuzulassen ist. Diese Anstze fhren zu der einfachen Regel: *Bei nderung des Bezugssystemes ist die vorausgesetzte Kraft derart als Kraft in den neuen Raumkoordinaten anzusetzen, da dabei der zugehrige Vektor mit den Komponenten*

$$\dot{t}X, \dot{t}Y, \dot{t}Z, \dot{t}T,$$

wo

$$T = \frac{1}{c^2} \left(\frac{\dot{x}}{\dot{t}} X + \frac{\dot{y}}{\dot{t}} Y + \frac{\dot{z}}{\dot{t}} Z \right)$$

die durch c^2 dividierte Arbeitsleistung der Kraft im Weltpunkte ist, sich unverndert erhlt. Dieser Vektor ist stets normal zum Bewegungsvektor in P . Ein solcher, zu einer Kraft in P gehrender Kraftvektor soll ein *bewegender Kraftvektor* in P heien.

*) M. Planck, Zur Dynamik bewegter Systeme, Sitzungsberichte der k. preuischen Akademie der Wissenschaften zu Berlin, 1907, S. 542 (auch Annalen der Physik, Bd. 26, 1908, S. 1).

**) H. Minkowski, Die Grundgleichungen fr die elektromagnetischen Vorgnge in bewegten Krpern, Nachrichten der k. Gesellschaft der Wissenschaft zu Gttingen, mathematisch-physikalische Klasse, 1908, S. 53 und Mathematische Annalen, Bd. 68, 1910, S. 527; diese Ges. Abhandlungen, Bd. II, S. 352.

Nun werde die durch P laufende Weltlinie von einem substantiellen Punkte mit konstanter *mechanischer Masse* m beschrieben. Das m -fache des Bewegungsvektors in P heie der *Impulsvektor* in P , das m -fache des Beschleunigungsvektors in P der *Kraftvektor der Bewegung* in P . Nach diesen Definitionen lautet das Gesetz dafr, wie die Bewegung eines Massenpunktes bei gegebenem bewegenden Kraftvektor statthat*):

Der Kraftvektor der Bewegung ist gleich dem bewegenden Kraftvektor.

Diese Aussage fat vier Gleichungen fr die Komponenten nach den vier Achsen zusammen, wobei die vierte, weil von vornherein beide genannten Vektoren normal zum Bewegungsvektor sind, sich als eine Folge der drei ersten ansehen lt. Nach der obigen Bedeutung von T stellt die vierte zweifellos den Energiesatz dar. Als *kinetische Energie* des Massenpunktes ist daher das c^2 -fache der *Komponente des Impulsvektors nach der t -Achse* zu definieren. Der Ausdruck hierfr ist

$$m c^2 \frac{dt}{d\tau} = m c^2 \left| \sqrt{1 - \frac{v^2}{c^2}} \right|,$$

d. i. nach Abzug der additiven Konstante $m c^2$ der Ausdruck $\frac{1}{2} m v^2$ der Newtonschen Mechanik bis auf Gren von der Ordnung $1/c^2$. Sehr anschaulich erscheint hierbei die *Abhngigkeit der Energie vom Bezugssysteme*. Da nun aber die t -Achse in die Richtung jedes zeitartigen Vektors gelegt werden kann, so enthlt andererseits der Energiesatz, fr jedes mgliche Bezugssystem gebildet, bereits das ganze System der Bewegungsgleichungen. Diese Tatsache behlt bei dem errterten Grenzübergang zu $c = \infty$ ihre Bedeutung auch fr den axiomatischen Aufbau der Newtonschen Mechanik und ist in solchem Sinne hier bereits von Herrn J. R. Schtz**) wahrgenommen worden.

Man kann von vornherein das Verhltnis von Lngeneinheit und Zeiteinheit derart festlegen, da die natrliche Geschwindigkeitsschranke $c = 1$ wird. Fhrt man dann noch $\sqrt{-1} \cdot t = s$ an Stelle von t ein, so wird der quadratische Differentialausdruck

$$d\tau^2 = -dx^2 - dy^2 - dz^2 - ds^2,$$

also vllig symmetrisch in x, y, z, s , und diese Symmetrie bertrgt sich auf ein jedes Gesetz, das dem Weltpostulate nicht widerspricht. Man kann danach das Wesen dieses Postulates mathematisch sehr prgnant in die mystische Formel kleiden:

$$3 \cdot 10^5 \text{ km} = \sqrt{-1} \text{ sek.}$$

*) H. Minkowski, diese Ges. Abhandlungen, Bd. II, S. 400. — Vgl. auch M. Planck, Verhandlungen der Physikalischen Gesellschaft, Bd. 4, 1906, S. 136.

**) J. R. Schtz, Das Prinzip der absoluten Erhaltung der Energie, Nachrichten der k. Gesellschaft der Wissenschaften zu Gttingen, mathematisch-physikalische Klasse, 1897, S. 110.

V.

Die durch das Weltpostulat geschaffenen Vorteile werden vielleicht durch nichts so schlagend belegt wie durch Angabe der von einer



Fig. 3.

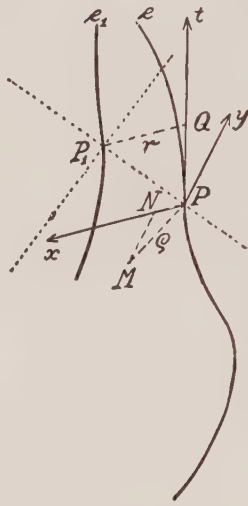


Fig. 4.

beliebig bewegten punktförmigen Ladung nach der Maxwell-Lorentzschen Theorie ausgehenden Wirkungen. Denken wir uns die Weltlinie eines solchen punktförmigen Elektrons mit der Ladung e und führen auf ihr die Eigenzeit τ ein von irgendeinem Anfangspunkte aus. Um das vom Elektron in einem beliebigen Weltpunkte P_1 veranlaßte Feld zu haben, konstruieren wir den zu P_1 gehörigen Vorkegel (Fig. 4). Dieser trifft die unbegrenzte Weltlinie des Elektrons, weil deren Richtungen überall die von zeitartigen Vektoren sind, offenbar in einem einzigen Punkte P . Wir

legen in P an die Weltlinie die Tangente und konstruieren durch P_1 die Normale P_1Q auf diese Tangente. Der Betrag von P_1Q sei r . Als der Betrag von PQ ist dann gemäß der Definition eines Vorkegels r/c zu rechnen. Nun stellt der Vektor in Richtung PQ vom Betrage e/r in seinen Komponenten nach den x -, y -, z -Achsen das mit c multiplizierte Vektorpotential, in der Komponente nach der t -Achse das skalare Potential des von e erregten Feldes für den Weltpunkt P_1 vor. Hierin liegen die von A. Liénard und von E. Wiechert aufgestellten Elementargesetze*).

Bei der Beschreibung des vom Elektron hervorgerufenen Feldes selbst tritt sodann hervor, daß die Scheidung des Feldes in elektrische und magnetische Kraft eine relative ist mit Rücksicht auf die zugrunde gelegte Zeitachse; am übersichtlichsten sind beide Kräfte zusammen zu beschreiben in einer gewissen, wenn auch nicht völligen Analogie zu einer Kraftschraube der Mechanik.

Ich will jetzt die von einer beliebig bewegten punktförmigen Ladung auf eine andere beliebig bewegte punktförmige Ladung ausgeübte ponderomotorische Wirkung beschreiben. Denken wir uns durch den Weltpunkt

*) A. Liénard, Champ électrique et magnétique produit par une charge concentrée en un point et animée d'un mouvement quelconque, L'Éclairage électrique, T. 16, 1898, pp. 5, 53, 106; E. Wiechert, Elektrodynamische Elementargesetze, Archives Néerlandaises des Sciences exactes et naturelles (2), T. 5, 1900, S. 549.

P_1 die Weltlinie eines zweiten punktförmigen Elektrons von der Ladung e_1 führend. Wir bestimmen P , Q , r wie vorhin, konstruieren sodann (Fig. 4) den Mittelpunkt M der Krümmungshyperbel in P , endlich die Normale MN von M aus auf eine durch P parallel zu QP_1 gedachte Gerade. Wir legen nun, mit P als Anfangspunkt, ein Bezugssystem folgendermaßen fest, die t -Achse in die Richtung PQ , die x -Achse in die Richtung QP_1 , die y -Achse in die Richtung MN , womit schließlich auch die Richtung der z -Achse als normal zu den t -, x -, y -Achsen bestimmt ist. Der Beschleunigungsvektor in P sei $\ddot{x}, \ddot{y}, \ddot{z}, \ddot{t}$, der Bewegungsvektor in P_1 sei $\dot{x}_1, \dot{y}_1, \dot{z}_1, \dot{t}_1$. Jetzt lautet der von dem ersten beliebig bewegten Elektron e auf das zweite beliebig bewegte Elektron e_1 in P_1 ausgeübte bewegendende Kraftvektor:

$$-ee_1\left(\dot{t}_1 - \frac{\dot{x}_1}{c}\right)\mathfrak{R},$$

wobei für die Komponenten $\mathfrak{R}_x, \mathfrak{R}_y, \mathfrak{R}_z, \mathfrak{R}_t$ des Vektors $\mathfrak{R}$ die drei Relationen bestehen:

$$c\mathfrak{R}_t - \mathfrak{R}_x = \frac{1}{r^2}, \quad \mathfrak{R}_y = \frac{\dot{y}}{c^2 r}, \quad \mathfrak{R}_z = 0$$

und viertens dieser Vektor $\mathfrak{R}$ normal zum Bewegungsvektor in P_1 ist und durch diesen Umstand allein in Abhängigkeit von dem letzteren Bewegungsvektor steht.

Vergleicht man mit dieser Aussage die bisherigen Formulierungen*) des nämlichen Elementargesetzes über die ponderomotorische Wirkung bewegter punktförmiger Ladungen aufeinander, so wird man nicht umhin können zuzugeben, daß die hier in Betracht kommenden Verhältnisse ihr inneres Wesen voller Einfachheit erst in vier Dimensionen enthüllen, auf einen von vornherein aufgezungenen dreidimensionalen Raum aber nur eine sehr verwickelte Projektion werfen.

In der dem Weltpostulate gemäß reformierten Mechanik fallen die Disharmonien, die zwischen der Newtonschen Mechanik und der modernen Elektrodynamik gestört haben, von selbst aus. Ich will noch die Stellung des *Newtonschen Attraktionsgesetzes* zu diesem Postulate berühren. Ich will annehmen, wenn zwei Massenpunkte m, m_1 ihre Weltlinien beschreiben, werde von m auf m_1 ein bewegendender Kraftvektor ausgeübt genau von dem soeben im Falle von Elektronen angegebenen Ausdruck, nur daß statt $-ee_1$ jetzt $+mm_1$ treten soll. Wir betrachten nun speziell den Fall, daß der Beschleunigungsvektor

*) K. Schwarzschild, Nachrichten der k. Gesellschaft der Wissenschaften zu Göttingen, mathematisch-physikalische Klasse, 1903, S. 132. — H. A. Lorentz, Enzyklopädie der mathematischen Wissenschaften, V, Art. 14, S. 199.

von m konstant Null ist, wobei wir dann t so einführen mögen, daß m als ruhend aufzufassen ist, und es erfolge die Bewegung von m_1 allein mit jenem von m herrührenden bewegenden Kraftvektor. Modifizieren wir nun diesen angegebenen Vektor zunächst durch Hinzusetzen des Faktors $t^{-1} = \sqrt{1 - \frac{v^2}{c^2}}$, der bis auf Größen von der Ordnung $1/c^2$ auf 1 hinauskommt, so zeigt sich*), daß für die Orte x_1, y_1, z_1 von m_1 und ihren zeitlichen Verlauf genau wieder die Keplerschen Gesetze hervorgehen würden, nur daß dabei an Stelle der Zeiten t_1 die Eigenzeiten τ_1 von m_1 eintreten würden. Auf Grund dieser einfachen Bemerkung läßt sich dann einsehen, daß das vorgeschlagene Anziehungsgesetz verknüpft mit der neuen Mechanik nicht weniger gut geeignet ist die astronomischen Beobachtungen zu erklären als das Newtonsche Anziehungsgesetz verknüpft mit der Newtonschen Mechanik.

Auch die Grundgleichungen für die elektromagnetischen Vorgänge in ponderablen Körpern fügen sich durchaus dem Weltpostulate. Sogar die von Lorentz gelehrte Ableitung dieser Gleichungen auf Grund von Vorstellungen der Elektronentheorie braucht zu dem Ende keineswegs verlassen zu werden, wie ich anderwärts zeigen werde**).

Die ausnahmslose Gültigkeit des Weltpostulates ist, so möchte ich glauben, der wahre Kern eines elektromagnetischen Weltbildes, der von Lorentz getroffen, von Einstein weiter herausgeschält, nachgerade vollends am Tage liegt. Bei der Fortbildung der mathematischen Konsequenzen werden genug Hinweise auf experimentelle Verifikationen des Postulates sich einfinden, um auch diejenigen, denen ein Aufgeben altgewohnter Anschauungen unsympathisch oder schmerzlich ist, durch den Gedanken an eine prästabilisierte Harmonie zwischen der reinen Mathematik und der Physik auszusöhnen.

*) H. Minkowski, diese Ges. Abhandlungen, S. 403.

**) Dieser Gedanke ist ausgeführt in der Arbeit: Eine Ableitung der Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern vom Standpunkte der Elektronentheorie. Aus dem Nachlaß von Hermann Minkowski bearbeitet von Max Born in Göttingen. Mathematische Annalen, Bd. 68, 1910, S. 526; diese Ges. Abhandlungen, Bd. II, S. 405.

REDE AUF DIRICHLET

XXXIII.

Peter Gustav Lejeune Dirichlet und seine Bedeutung für die heutige Mathematik.

(Rede, gehalten in der Festsitzung der Göttinger Mathematischen Gesellschaft
am 13. Februar 1905.)

(Jahresbericht der Deutschen Mathematiker-Vereinigung, Band 14, S. 149—163.)*

Hochgeehrte Anwesende!

Gedanken an Zeit und Raum sind des Mathematikers ständige Gefährten, die er aber von sich weist, wenn er das Reich der reinen Zahl betritt. Heute werden wir an Herrlichkeiten tief im Innern des Zahlenreiches herantreten, und die weihevollte Stimmung des Augenblicks, der Einfluß des Ortes gerade sind es, die unsere Schritte dorthin lenken.

Wir begehen die hundertjährige Wiederkehr des Geburtstages von Lejeune Dirichlet, und wir denken zuerst an die Umstände, die uns das Recht geben, Dirichlet als zu Göttingen gehörig zu betrachten.

Am 23. Februar 1855 war Gauß gestorben. Der damalige Kurator von Warnstedt wandte sich an Wilhelm Weber als das zunächst kompetente Mitglied der Honorenfakultät mit dem Ersuchen, ihm Bericht zu erstatten, wie die im Lehrkörper der Universität entstandene Lücke so würdig als möglich ausgefüllt werden könnte. Für die Entschließungen leitend, hieß es in dem Schreiben, werde vor allem die Rücksicht sein, welche bisher in ähnlichen Lagen der Universität festgehalten sei. Es ist das der Gesichtspunkt, daß die Universität nicht bloß eine hohe Schule für den Unterricht der Studierenden, nicht bloß eine Bewahrerin bereits erworbener wissenschaftlicher Kenntnisse sein, sondern auch an dem Fortbau der Wissenschaft als eines Gemeinguts der Menschheit arbeiten soll. Diesem als leitend festgehaltenen Gesichtspunkt verdankt Göttingen die ihm eigentümliche Weltstellung. Im Gebiete der höheren Mathematik und der auf die Mathematik begründeten naturwissenschaftlichen Unter-

*) Der Rede war im Sonderabdruck ein Bildnis Lejeune Dirichlets als Titelbild beigegeben. (Anm. d. Herausg.)

suchungen hat die Georgia Augusta vielleicht diesen Rang am glänzendsten bewährt. Universitäten, welche eine wesentliche und bedeutsame Stellung in dem geistigen Leben der Nation einnehmen, sind Individuen mit einem geschichtlich festgestellten Charakter, in dessen Bewahrung die Gewähr ihrer Kraft und ihres Segens für das Ganze ruht.

Weber sandte ein eingehendes Gutachten, er verbreitete sich ausführlich über das Verhältnis der Mathematik zu den ihr verwandten Naturwissenschaften, denn Gauß hatte seine Wirksamkeit zugleich über alle von der Mathematik beherrschten Wissenschaften, über die Physik ebenso wie über die Astronomie, ausgedehnt. Weber kam zu dem Schlusse, daß als der würdigste Nachfolger von Gauß wohl einstimmig, in und außer Deutschland, Lejeune Dirichlet, damals in Berlin tätig, betrachtet werden möchte. Es wurden auf der Stelle Verhandlungen mit Dirichlet eingeleitet. Dieser selbst legte auf Beschleunigung der Vokation Wert, weil sonst alle die Versuche, ihn in Berlin zu halten, ihm die Annahme des Rufes sehr erschweren würden. Schon nach kurzer Frist konnte Warnstedt an den Minister berichten: Dirichlet kommt gerne nach Göttingen, weil er es für den höchsten Ruhm hält, Gauß' Nachfolger zu werden. Es ist das eine wahrhaft glanzvolle Akquisition, die schönste, die Göttingen in diesem Fache hätte machen können.

Der Universität Göttingen, welche ein halbes Jahrhundert hindurch den Ruhm genossen hatte, den ersten aller lebenden Mathematiker zu besitzen, war es gelungen, sich diesen Ruhm auch ferner zu erhalten.

Einige Daten aus dem Leben Dirichlets werden Ihr Interesse beanspruchen dürfen; sie sind zumeist der ausgezeichneten Gedächtnisrede entnommen, die Kummer in der Berliner Akademie seinem um fünf Jahre älteren Freunde gewidmet hat.

Peter Gustav Lejeune Dirichlet wurde in Düren bei Aachen geboren, heute vor hundert Jahren. Auf dem Gymnasium in Köln hatte er zu seinem Lehrer in der Mathematik den nachmals in der Elektrodynamik berühmt gewordenen Georg Simon Ohm. Sechzehn Jahre alt, kehrte er mit dem Abgangszeugnis für die Universität nach Hause zurück und beredete die anfänglich widerstrebenden Eltern, seinem entschiedenen Wunsche nachzugeben, ihn Mathematik studieren zu lassen.

Das mathematische Studium auf den preußischen und auf den übrigen deutschen Universitäten lag damals arg darnieder. Die Vorlesungen erhoben sich nur wenig über das Gebiet der Elementarmathematik. In Frankreich dagegen stand die Mathematik in voller Blüte. Laplace, Legendre, Fourier, Poisson, Cauchy wirkten zusammen, Paris zu einem glänzenden Sitze der mathematischen Wissenschaft zu machen.

Dirichlet wandte sich nach Paris, er verblieb dort viereinhalb Jahre,

den größten Teil der Zeit als Lehrer im Hause des Generals Foy tätig, der eine hervorragende Persönlichkeit war und als Politiker eine bedeutende Rolle gespielt hat. Mehr noch als das Hören der Vorlesungen wurde in dieser Zeit für Dirichlets spätere wissenschaftliche Richtung bestimmend das Studium der Gaußischen *Disquisitiones Arithmeticae*. Dirichlet hat dieses Werk nicht nur einmal oder mehrere Male durchstudiert, sein ganzes Leben hindurch hat er nicht aufgehört, die Fülle der tiefen Gedanken, die es enthält, sich immer wieder zu vergegenwärtigen. Sartorius von Waltershausen erzählte einmal: So wie gewisse Geistliche mit ihrem Gebetbuch umherziehen, pflegt Dirichlet nur in Begleitung eines ganz verlesenen, aus dem Einband gewichenen Exemplars der *Disquisitiones Arithmeticae* auf alle seine Reisen zu gehen.

Durch eine Arbeit zum Fermatschen Satze erweckte Dirichlet die Aufmerksamkeit der Pariser Akademie. Infolge dieses glänzenden Debuts kam Dirichlet namentlich mit Fourier in nähere Verbindung. Fourier fand in ihm einen jungen Mann, dem er sein mathematisches Herz ganz eröffnen konnte und von dem er nicht bloß bewundert, sondern auch verstanden wurde. Durch Fourier wurde Alexander von Humboldt auf Dirichlet aufmerksam gemacht. Humboldt glaubte, daß unter den damaligen Verhältnissen Frankreichs der Aufschwung der Wissenschaften gehemmt und auch für große Talente in Paris wenig Aussicht wäre, und er brachte in Dirichlet den schon erwogenen Entschluß zur Rückkehr ins Vaterland zur Reife. Durch Verwendung Humboldts, dessen Bemühungen auch Gauß unterstützte, erhielt Dirichlet 1827 eine Anstellung in Breslau. Auf der Reise dorthin wählte er den Weg über Göttingen, um Gauß persönlich kennen zu lernen.

In Breslau blieb Dirichlet wenig über ein Jahr, dann wurde er nach Berlin gerufen, wo er zuerst einen Lehrauftrag an der Kriegsschule erhielt und bald für die dortige Universität und die Akademie gewonnen wurde. Dirichlet wirkte an der Berliner Hochschule in ausgezeichneter Weise 27 Jahre hindurch, davon sieben Jahre im Verein mit seinem Freunde Jacobi. Eine Berufung nach Heidelberg im Jahre 1846 hat er abgelehnt.

Im Berichte Webers findet sich noch eine Tatsache, die nicht bekannt geworden ist. Schon im Jahre 1828 hegte Gauß im Interesse der Universität und in der Hoffnung des gemeinsamen wissenschaftlichen Zusammenwirkens den lebhaftesten Wunsch und sprach ihn gelegentlich gegen Humboldt aus, daß Dirichlet nach Göttingen berufen werden möge. Nach Thibauts Tode 1832 wollte Gauß selbst einen entsprechenden Antrag bei dem Kuratorium stellen, er nahm davon nur Abstand, weil er sich sagen mußte, Dirichlet würde Berlin nicht verlassen haben

infolge der soeben erst durch seine Verheiratung eingegangenen Verbindung mit dem Mendelssohnschen Hause, welches damals in Berlin den ausgezeichnetsten Vereinigungspunkt für Kunst und Wissenschaft bildete. Dirichlet hatte sich im Jahre 1831 mit Rebecca Mendelssohn-Bartholdy, der jüngeren Schwester von Felix Mendelssohn und von Fanny Hensel, vermählt.

„Im Herbst 1855 siedelte Dirichlet nach Göttingen über. Er richtete sich hier in einem eigenen, angenehm gelegenen Hause mit Garten (Mühlenstraße 1) ganz nach seinem Gefallen ein, und die Ruhe der kleineren Stadt, welche er seit seiner Jugend nicht mehr genossen hatte, ersetzte ihm hinreichend die äußeren Annehmlichkeiten des großstädtischen Lebens in Berlin.“ An der Universität fand er zwar nicht einen so großen Kreis von Zuhörern, als er in Berlin verlassen hatte, allein sein Ruf als Lehrer war nicht minder anerkannt als sein wissenschaftlicher Ruf und zog viele junge Mathematiker nach Göttingen.

Es sollte ihm jedoch nur noch eine kurze Wirksamkeit gegönnt sein.

Im Sommer des Jahres 1858, nach dem Schlusse seiner Vorlesungen, reiste er nach der Schweiz und hielt sich in Montreux auf, weniger zu seiner Erholung, als mit dem Vorhaben, daselbst eine in der Göttinger Sozietät zu haltende Gedächtnisrede auf Gauß auszuarbeiten. Dort wurde er plötzlich von einer akuten Herzkrankheit ergriffen und kehrte todkrank nach Göttingen zurück. Die augenblickliche Lebensgefahr konnte zwar abgewandt werden, da traf Dirichlet der plötzliche Verlust seiner Frau. Dieser Schlag wendete seine Krankheit wieder zum Schlimmeren und nach schweren Leiden erlag er derselben am 5. Mai 1859, nur 54 Jahre alt, in der besten Schaffenskraft, in frischem Besitze wunderbarer Geheimnisse, die er mit sich hinübernahm. Er ruht auf dem alten Kirchhofe, der an der Weender Chaussee liegt; das Grab, das ihn und seine Frau aufnahm, befindet sich wenige Schritte rechts hinter dem Bürgerdenkmal und ist leicht kenntlich an einer hohen steinernen Umfriedung.

Seit jener Zeit ist fast ein halbes Jahrhundert weiterer intensivster Entwicklung der mathematischen Wissenschaft vorübergezogen. In allen ihren Teilen sehen wir noch die Impulse des Dirichletschen Geistes nachwirken. Suchen wir das Lebenswerk Dirichlets als Ganzes zu beurteilen, so ist es ein ununterbrochenes Zeugnis für die Verbrüderung der mathematischen Disziplinen, für die Einheit unserer Wissenschaft. Unablässig war sein Sinn darauf gerichtet, zwischen getrennt bestehenden Gedankensphären die Brücke zu schlagen und unabhängig gezeigte Erfolge zu fruchtbarer Wechselwirkung zu verschmelzen.

Das Feld, auf welchem unstreitig die bewundernswürdigsten Offenbarungen Dirichlets liegen, ist die höhere Arithmetik, die Lehre von den Eigenschaften der ganzen Zahlen. Es kann nur derjenige Dirichlet in richtigem Lichte erfassen, der die Eigenart dieses Zweiges der Mathematik zu schätzen versteht.

Ich muß davon etwas ausführlicher sprechen, wenn ich Ihr Interesse für Dirichlets größte Leistungen wachrufen will. Freilich, die Arithmetik von heute ist nicht ganz die Arithmetik aus den Zeiten der Disquisitiones. An die Stelle der erfindungsreichen schöpferischen Phantasie eines Gauß und Dirichlet, die ein verheißenes Land suchte, ist mehr der systematische Anbau eines sicher erworbenen Bodens getreten. Für Gauß war die Arithmetik die Königin der Mathematik, sie ist es auch heute noch. Aber ihr Wesen ist selbst vielen Mathematikern unnahbar, man hört von der fortschreitenden Arithmetisierung *aller* mathematischen Wissenszweige sprechen, und manche halten deshalb die Arithmetik nur noch für eine zweckmäßige Staatsverfassung, die sich das ausgedehnte Reich der Mathematik gibt. Ja, zuletzt werden einige in ihr nur noch die hohe Polizei sehen, welche befugt ist, auf alle verbotenen Vorgänge im weitverzweigten Gemeinwesen der Größen und Funktionen zu achten. Die Arithmetik steht einzig da durch „die Einfachheit ihrer Grundlagen, die Genauigkeit ihrer Begriffe, die Reinheit ihrer Wahrheiten“, diese Ruhmestitel ihrer Unbestechlichkeit werden ihr von niemandem aberkannt.

Erwecken wir für einen Moment die alte Arithmetik aus ihrem Dornröschenschlaf.

Fast alle, die sich ernstlich um die Arithmetik bemüht haben, gaben sich ihr bald mit einer gewissen Leidenschaft hin. Mit einer leichten Variation des Ausspruchs, den Novalis auf die Mathematiker im allgemeinen geprägt hat, darf füglich behauptet werden: Der echte Arithmetiker ist Enthusiast per se. Ohne Enthusiasmus keine Arithmetik.

Allerdings sind unter den Mathematikern selbst manche anzutreffen, denen der Reiz für dieses Gebiet vollkommen abzugehen scheint. Vielleicht liegt die Ursache davon in dem hohen Grade von Abstraktionsfähigkeit, welchen die arithmetischen Begriffe zu ihrer völligen Beherrschung erfordern. Vielleicht auch rührt jene Abneigung von einer Ungeduld her, die der mathematischen Gedankenarbeit erst im Momente unmittelbarer praktischer Verwendbarkeit ihr Interesse zuwenden mag. In letzterer Hinsicht bin ich übrigens für die Zahlentheorie Optimist und hege still die Hoffnung, daß wir vielleicht gar nicht weit von dem Zeitpunkt entfernt sind, wo die unverfälschteste Arithmetik gleichfalls in Physik und Chemie Triumphe feiern wird, und sagen wir z. B., wo wesentliche Eigenschaften der Materie als mit der Zerlegung der Primzahlen in

zwei Quadrate im Zusammenhang stehend erkannt werden. An jenem Tage werden den Arithmetikern von allen Seiten Huldigungen dargebracht werden. Einstweilen aber würden in einen Briefsteller für Arithmetiker noch zweckmäßig alle die resignierten Äußerungen hineingehören, die Gauß und Dirichlet sich über die geringe Verbreitung zahlentheoretischen Sinnes schrieben.

Schon jetzt nehmen wir in der Zahlentheorie mehrere, wenn auch vielleicht nur äußerliche Analogien mit der Physik wahr. Durch die unendliche Reihe der Zahlen bietet sich die weiteste Möglichkeit des Experimentierens. Jede Zahl ist wie ein Individuum für sich, das in der mannigfaltigsten Weise die Eigenschaften großer Klassen in sich vereinigt. Jede Zahl ist daher ein Objekt für Versuche zur Entdeckung allgemeiner Gesetze. Dem Arithmetiker werden in seiner Beschäftigung die Ziffern das, was die Farben dem bildenden Künstler sind. Er mischt dieses Handwerkszeug, um wunderbare Geschehnisse, eindrucksvolle Charaktere festzuhalten, indes ein anderer Mathematiker höchstens mit rohen Strichen den Effekt von Groß und Klein hervorzurufen, die Distanz der Fixsterne neben dem Durchmesser eines Moleküls zu malen versteht.

Auch, daß auffällige Zusammenhänge oft früher wahrgenommen werden, als die inneren Gründe dafür offenbar liegen, ist ganz ein Charakteristikum einer Erfahrungswissenschaft. Manche derartige Erscheinungen in der Zahlentheorie mögen in den Folgen, wie sie den Weg zeigten, um in neue unbekannte Tiefen der Wissenschaft einzudringen, wohl mit der Feststellung von Weber und Kohlrausch zu vergleichen sein, daß der Quotient der elektromagnetischen und der elektrostatischen Einheit der Elektrizitätsmenge genau mit der Größe der Lichtgeschwindigkeit übereinstimmt. So bildete ein merkwürdiges Endergebnis, auf das Dirichlet bei der Berechnung der Klassenanzahl der quadratischen Formen mit komplexen Koeffizienten geführt wurde, für einen mathematischen Forscher der Gegenwart das Einfallstor, durch welches er in das Reich der Zahlentheorie eindrang, um dort neue blühende Provinzen zu erschließen.

Andererseits finden wir in der Geschichte der Arithmetik, daß oft bedeutende Fortschritte an ganz unscheinbare Phänomene anknüpfen. In der Theorie der Kreisteilung war Gauß darauf geführt worden, gewisse Ausdrücke zu betrachten, — jetzt nennt man sie kurzweg die Gaußischen Summen, — sie definieren im wesentlichen in einem regulären Polygon von n Seiten einen gewissen ausgezeichneten Punkt. Es ergab sich, daß der Punkt auf einer bestimmten Geraden durch den Mittelpunkt in bestimmter Entfernung von letzterem liegen müsse, es blieb aber zweifelhaft, ob rechts oder links. In jedem Versuchsfalle ergab sich die Lage rechts, aber der Beweis dieser so einfach klingenden Tatsache bot die

allergrößten Schwierigkeiten dar und führte Gauß tief in das Formelgebiet der elliptischen Funktionen hinein. Dirichlet hat später auf einem sehr sinnreichen neuen Wege, durch Benutzung der Fourierschen Reihen und eines Integrals, das in der Fresnelschen Theorie der Beugung des Lichtes eine Rolle spielt, die Gaußischen Summen mitsamt dem zweifelhaften Vorzeichen zu berechnen gelehrt.

Ich habe etwas ausführlicher von dem Wesen der Zahlentheorie gesprochen. Nun will ich auf einzelne Funde von Dirichlet eingehen, wie sie sich in den heutigen Stand der mathematischen Wissenschaft einordnen. Der Brennpunkt der heutigen Zahlentheorie, von dem aus auch ihre funktionentheoretischen Beziehungen ausstrahlen, ist die Theorie der algebraischen Zahlkörper. Diese Theorie ruht auf drei Grundpfeilern, dem Satze von der eindeutigen Zerlegbarkeit in Primideale, dem Satze von der Existenz der Einheiten und dem Satze von der transzendenten Bestimmung der Klassenanzahl.*) Von diesen Pfeilern hat die beiden letzten Dirichlet allein errichtet, den zuerst genannten hat Kummer entworfen und es haben ihn dann Dedekind und Kronecker, Schüler Dirichlets, in unabhängiger Arbeit ausgebaut, und jener hat ihn sogleich jedermann sichtbar aufgeführt, dieser ihn lange hindurch verhüllt gehalten. Nun einiges Nähere zum Verständnis jener glänzendsten Entdeckungen Dirichlets.

Selbst die in der Zahlenlehre wenig Bewanderten wissen von der Gaußischen Abhandlung, in der die geometrische Deutung der komplexen Größen gelehrt wurde. In jener Abhandlung hat Gauß das Feld der Zahlentheorie außerordentlich erweitert. Die konsequente Fortentwicklung des Gaußischen Gedankens hat dazu geführt, dem Begriffe der ganzen Zahl eine noch unendlich mannigfaltigere Ausbildung zu geben. Während die früheren ganzen Zahlen, nun zur Unterscheidung ganze *rationale* genannt, sich additiv und subtraktiv aus 1 zusammensetzen, haben die neuen ganzen Zahlen, die *algebraischen* ganzen Zahlen genannt, die charakteristische Eigenschaft, daß irgendeine Potenz von ihnen sich additiv und subtraktiv aus früheren Potenzen zusammensetzen läßt. Derartige Mengen von Zahlen, welche auseinander ausschließlich durch Anwendung der vier Spezies hervorgehen, werden zu den sogenannten *Zahlkörpern* zusammengefaßt. Auch in diesen erweiterten Zahlbereichen blieb der multiplikative Aufbau aller ganzen Zahlen aus möglichst einfachen Primelementen das grundlegende Problem. Vor allem erhebt sich da die Frage nach allen

*) Vgl. Hilbert, Vorwort zu seinem Bericht über die Theorie der algebraischen Zahlkörper in Bd. 4 der Jahresberichte der Deutschen Mathematiker-Vereinigung.

denjenigen ganzen Zahlen, welche auch in der *Eins* als Faktor enthalten sein können, wie dieses von den rationalen ganzen Zahlen nur ausschließlich die zwei Werte 1 und -1 tun. Diese Zahlen eines Zahlkörpers werden *Einheiten* genannt, sie durchdringen die übrigen Zahlen und beeinflussen deren innerste Natur, ich möchte sagen, wie der Lichtäther die Erscheinungen der Materie.

Vor Dirichlet war die Frage nach den Einheiten nur für den einfachsten Fall der reellen Körper, welche von einer quadratischen Gleichung abhängen, erledigt. Aber indem hier der *Nachweis der Existenz* der Einheit gleichzeitig mit einem speziellen *Verfahren zur Auffindung* der Einheit Hand in Hand ging, konnte man versucht sein, die Wichtigkeit des besonderen Algorithmus zu überschätzen, und hätte alle Kräfte vergeuden können in dem Unternehmen, das angetroffene formale Hilfsmittel auf allgemeinere Fälle zu übertragen. Hier hat nun Dirichlet mit der ganzen Schärfe vorurteilsfreier Auffassung eingesetzt, den tiefliegenden Kern des Problems herausgeschält und den ganzen merkwürdigen Organismus der Einheiten in vollster Klarheit auseinandergelegt. Er zeigte, daß in der Regel in einem Zahlkörper *unendlich viele* Einheiten vorhanden sind, und daß sie sich alle aus einer *endlichen* Anzahl unter ihnen durch Multiplikation und Potenzierung ableiten lassen. Es ist wahrhaft bewundernswürdig, in welche einfache Form Dirichlet zuletzt den Beweis seines gewaltigen Theorems zu gießen verstand. Fast sieht es aus, als ob als wesentlichstes Hilfsmittel nur ein ganz alltägliches Prinzip verbleibt: Tut man in eine Anzahl von Schubfächern eine größere Anzahl von Gegenständen, als man Schubfächer hat, so finden sich notwendig in wenigstens einem Fache gleichzeitig zwei oder mehrere Gegenstände vor. Allerdings werden die Dirichletschen Schubfächer Größenbereiche und die darin aufgehobenen Gegenstände natürliche Logarithmen.

Es wird erzählt, daß nach langjährigen vergeblichen Bemühungen um das schwierige Problem Dirichlet die Lösung in Rom in der Sixtinischen Kapelle während des Anhörens der Ostermusik ergründet hat. Inwieweit dieses Faktum für die von manchen behauptete Wahlverwandtschaft zwischen Mathematik und Musik spricht, wage ich nicht zu erörtern.

Die zweite große zahlentheoretische Entdeckung Dirichlets, die Formel für die Klassenanzahl, lag in einem Gebiete, das vollkommen unabhängig von der Theorie der algebraischen Zahlkörper aufwuchs, in der Theorie der quadratischen Formen. Es sind aber inzwischen die Zweige der Arithmetik so dicht zusammengewachsen, daß ich auch die Bedeutung dieser anderen Dirichletschen Großtat auseinandersetzen kann, während wir in dem zuletzt berührten Ideenkreise verbleiben.

In den höheren Zahlkörpern hörte höchst merkwürdigerweise die Eindeutigkeit der Zerfällung der ganzen Zahlen in unzerlegbare Primelemente auf. Fast schien es, als wenn an dieser Schwierigkeit überhaupt die Weiterführung der Theorie der Körper in der einmal eingeschlagenen Richtung scheitern müßte. Aus dem Labyrinth, in das die Arithmetik hier verstrickt erschien, fand Kummer durch einen gänzlich neuen Gedanken den rettenden Ausweg. Die nicht weiter zerlegbaren Zahlen durften eben noch nicht als die letzten Urelemente angesehen werden, auch sie waren das Produkt einer weiter reichenden Zusammensetzung, die wahren letzten Primelemente freilich ließen sich nicht mehr als reine Zahlen ab scheiden. Aber es ergab sich ein Hilfsmittel, das Vorhandensein eines bestimmten Primelementes unzweifelhaft nachzuweisen, und der fundamentale Satz von der eindeutigen Zerlegbarkeit in Faktoren blieb gerettet.

Kummer nannte die neuen Bildungen, über deren Vorhandensein als Faktoren in jedem einzelnen Falle mit Sicherheit entschieden werden konnte, ohne daß eine tatsächliche Division ausführbar war, *ideale Zahlen*. Dedekind gestaltete in der Folge diesen Begriff noch durchsichtiger. Eine bestimmte ideale Zahl ist völlig charakterisiert durch die Gesamtheit aller wirklichen Zahlen, in denen sie sich als Faktor zeigt, diese Gesamtheit ist geradezu ihr Abbild, und bei dieser Auffassung ersetzte Dedekind die Kummersche Bezeichnung ideale Zahl durch den kürzeren Namen *Ideal*. So ist es gekommen, daß die Mathematiker auch solche Ideale besitzen, die ausschließlich durch die Vernunft geboten sind; es sind das nicht die einzigen Ideale des Mathematikers, aber auch sie sind derart, daß derjenige, der sie kennt, sich an ihnen begeistert.

Nun vermögen sich zwei verschiedene Ideale in gewisser Weise auszutauschen, ein Vorgang, der durch die Analogie mit den chemischen Umwandlungen sehr einleuchtend sein wird. Man fügt zu den sämtlichen Zahlen, die ein gegebenes Ideal enthalten, einen wirklichen Faktor hinzu; es läßt sich aus den zusammengesetzten Produkten ein anderer wirklicher Faktor abspalten, und nun bleiben genau die sämtlichen Zahlen übrig, die ein bestimmtes zweites Ideal enthalten. In solchem Falle heißen die betreffenden zwei Ideale einander *äquivalent* und sie werden zu einer *Klasse* gerechnet. In denjenigen besonderen Fällen, wo der ursprüngliche Begriff der Primzahlen in Kraft bleibt, sind alle Ideale gegenseitig austauschbar und gibt es daher nur eine einzige Idealklasse. In *allen* Fällen aber erweist sich für einen Zahlkörper die Anzahl der sämtlichen vorhandenen Idealklassen als eine *endliche*, und diese *Anzahl der Idealklassen* ist das bedeutungsvollste Attribut eines algebraischen Zahlkörpers. Dirichlet nun ist es zuerst gelungen, das große Rätsel der Klassenanzahl durch ihre Bestimmung auf transzendente Wege zu lösen.

Dirichlet behandelte nur die quadratischen Zahlkörper, aber seine Methode ist in der Folge für die Erledigung der allgemeinsten Fälle zu reichend geblieben.

In dem Nachlasse von Gauß fand sich ein Fragment einer Abhandlung, die der Sozietät überreicht werden sollte. Sie beginnt: Sechszund-dreißig Jahre sind schon verflossen, seit ich die Prinzipien des wunderbaren Zusammenhangs, welchem die vorliegende Arbeit gewidmet ist, entdeckte. Aus diesem Fragmente geht unzweifelhaft hervor, daß Gauß, wie er in so vielen Richtungen spätere Funde vorweggenommen hat, so auch den Ausdruck für die Klassenanzahl schon 1801 gekannt hat. Aber der Weg, den Gauß bei dessen Herleitung einschlug, bricht in dem Fragment mit der Schwierigkeit eines Konvergenznachweises ab, dessen vermutliche Erledigung zu rekonstruieren nicht gelungen ist. Das Gaußsche Fragment handelt in seinem ausgeführten Teile von der Bestimmung des Flächeninhalts einer Figur, speziell des Kreises. Sie sehen daraus, wie die höchsten Probleme, welche die Arithmetik darbietet, immer wieder zu neuer Durchforschung elementarer Begriffe, die man schon längst gesichert wähnen würde, die Veranlassung werden.

Daß Dirichlet die fragliche Schwierigkeit überhaupt nicht antraf, liegt daran, daß er von einem gänzlich neuen Ausdruck für den Flächeninhalt einer Figur ausgehen konnte. Die gewöhnliche, ich möchte sagen, *mikroskopische* Bestimmung eines Flächeninhalts, welche auch Gauß auseinandersetzt, besteht darin, auf die zu untersuchende Figur Quadratnetze mit immer engeren und engeren Maschen zu legen und die in die Figur fallenden Maschen zu zählen. Dirichlet denkt sich ein für allemal ein bestimmtes, unendliches Quadratnetz fest über die Ebene verbreitet. Die zu untersuchende Figur wird nun von einem festen Anfangspunkte aus kontinuierlich in allen Dimensionen gleichmäßig vergrößert, und für jeden Kreuzungspunkt des Netzes wird angemerkt, bei welchem Vergrößerungsverhältnis er gerade auf den Rand der Figur treten würde. Die Summe der reziproken Werte aller so bestimmten Vergrößerungszahlen für alle vorhandenen Netzpunkte würde noch unendlich sein; man summiere nun aber bestimmte gleich hohe, etwa $2s^{\text{te}}$ Potenzen aller dieser reziproken Werte, so kommt man auf eine durch die ursprüngliche Figur völlig charakterisierte Funktion $\zeta(s)$; diese gestattet eine analytische Fortsetzung für alle komplexen s und weist dann insbesondere für $s = 1$ einen einfachen Pol auf, dessen Cauchysches Residuum genau der Flächeninhalt der Figur wird. Ich möchte vorschlagen, zur leichteren Einbürgerung dieser fundamentalen Überlegung in der Analysis die eben beschriebene Formel den Dirichletschen *makroskopischen* Ausdruck eines Flächeninhalts zu nennen.

Daß Dirichlet auf diesen merkwürdigen Ausdruck überhaupt verfiel, hatte er dem glücklichen Umstande zu verdanken, daß dergleichen Summen aus gleich hohen Potenzen positiver abnehmender Größen, die man heutzutage allgemein als Dirichletsche Reihen bezeichnet, sich ihm bei einer früheren Gelegenheit ganz von selbst dargeboten hatten. Es handelte sich damals darum, einen Beweis für den Satz zu finden, daß jede arithmetische Progression von Zahlen, bei welcher nicht alle Individuen einen gemeinschaftlichen Faktor haben, stets auch unendlich viele Primzahlen enthält. Legendre hatte richtig erkannt, daß dieser Satz durch seinen elementaren Charakter sich als ein äußerst einschneidendes Reagensmittel bei den mannigfachsten arithmetischen Überlegungen darstellt. Aber der Versuch Legendres, einen Beweis des Satzes zu erbringen, scheiterte kläglich, darf man mit gutem Grunde sagen.

Es ist nun ein weiterer Ruhmestitel von Dirichlet, zuerst jenes fundamentale Theorem über die Primzahlen in arithmetischen Progressionen bewiesen zu haben. Die Idee des Beweises war ihm durch Reflexion über die Art erwachsen, wie Euler aus der Verwandlung der Summe der reziproken Werte der natürlichen ganzen Zahlen in ein nur von der Reihe der Primzahlen abhängendes Produkt geschlossen hatte, daß die Anzahl der Primzahlen notwendig eine unendliche ist. Die größte Schwierigkeit, welche Dirichlet an dieser Stelle entgegentrat, bestand in der Notwendigkeit, das Nichtverschwinden des Wertes einer gewissen Summe einzusehen; und gerade diese Schwierigkeit löste sich in höchst überraschender Weise und wurde der Anlaß zu der vorhin geschilderten Entdeckung. Nämlich jene Summe erwies sich als eine Klassenanzahl quadratischer Formen, und als Anzahl mußte sie notwendig mindestens 1, also von Null verschieden sein. Ich kann hier vielleicht hinzufügen, daß vor wenigen Jahren Mertens in Wien das Nichtverschwinden der fraglichen Summe auf einem direkten Wege durch Größenabschätzung in betreff der Glieder darzutun vermochte. Aber diese kühne Methode von Mertens ist, als wenn man durch Gestrüpp auf steilstem Wege zu einer Bergspitze empor klimmt, anstatt durch eine Landschaft mit herrlichen Ausblicken sanft anzusteigen.

Nur ein Teil der zahlentheoretischen Arbeiten von Dirichlet ist hier flüchtig berührt worden. Mit besonderer Hingabe war Dirichlet unablässig bemüht, die schwer zugänglichen Ideen von Gauß den Mathematikern näherzubringen, die starren Beweise desselben in flüssige und durchsichtige Methoden umzuwandeln. Von mancher der in diesem Sinne unternommenen Arbeiten, wie z. B. seiner geometrischen Theorie der ternären quadratischen Formen, gingen in der Folge starke Anregungen aus.

Kürzer wollen wir der zweiten Hauptrichtung folgen, die uns in den publizierten Arbeiten von Dirichlet hervortritt, und auf seine Beiträge zur mathematischen Physik eingehen. Diese zwei Richtungen, die Zahlentheorie und die mathematische Physik, so divergierend sie scheinen mögen, sind bei Dirichlet harmonisch vereinigt durch das Band der Integralrechnung. Wie man oft einen Maler aus jedem seiner Werke auf den ersten Blick an eigentümlichen Farbeffekten oder häufiger wiederkehrenden Stimmungen erkennt, so ist es eine anziehende und stets erfolgreiche Handhabung der Integrale, welche vielen Werken Dirichlets ein charakteristisches Gepräge verleiht.

Zu den Entdeckungen Dirichlets, die, ich möchte sagen, am populärsten geworden sind, gehören seine Methoden und Resultate im Gebiete der Fourierschen Reihen; sie sind bahnbrechend geworden für die moderne Behandlung der reellen Funktionen und in ihrem Werte geradezu unschätzbar wegen des Anstoßes, den sie zur Ausbildung der Mengenlehre gegeben haben. Die den Mathematikern sehr geläufige Geschichte der trigonometrischen Reihen ist ein fortwährender Kampf um die Beseitigung von Vorurteilen. Nur unter lebhaftem Widerstreit der Meinungen hatte die Tatsache Anerkennung gefunden, daß willkürliche Funktionen sich in nach den Sinus und Kosinus der Vielfachen des Arguments fortschreitende Reihen entwickeln lassen. Aber ein strenger Beweis für die Konvergenz und damit die Zulässigkeit der Reihen war noch nicht erbracht worden. Cauchy zog Hilfsmittel heran, die im Falle analytischer Funktionen zu partiellen Erfolgen führten, und vermochte von diesen Mitteln nicht abzugehen. Es erging dem großen Analytiker hier etwa wie den Brüdern Montgolfier, die nach der ersten geglückten Auffahrt ihrer Luftballons sich von dem dabei verwandten Heizmaterial nicht trennen konnten, weil sie auf Erzeugung elektrischen Rauches bedacht waren, wo es auf Verdünnung der Luft ankam. Dirichlet aber erkannte die wahren Bedingungen für den erstrebten Aufstieg zur vollen Höhe der willkürlichen Funktionen.

Dirichlets durch wunderbare Klarheit und Einfachheit berühmter Beweis, daß jedem Minimum potentieller Energie Stabilität des Gleichgewichts entspricht, ist eine zweite Tat solcher Art, daß er sich von verkehrten Hilfsmitteln freimachte, indem er statt der analytischen Regeln für die Bestimmung der Minima einer Funktion nur den ursprünglichen Begriff des Minimums heranzog. Immer wieder sind doch bedeutende Fortschritte in der Mathematik auf der Stelle errungen, sowie man nur wahrnimmt, daß Umstände, die stets als zusammengehörig betrachtet wurden, nichts miteinander zu tun haben. Hierin mag auch ein wesentlicher Grund liegen, weshalb so oft Mathematiker schon in jungen Jahren ganz unerwartete Erfolge davontreiben. Sie blicken in vielen Dingen

weniger voreingenommen. Es trägt jeder mathematische Soldat den Marschallstab im Tornister, wenn er nicht aus purer Disziplin auf alles Vorhandene schwört.

Leicht und von Grund aus zerstörte Dirichlet die naive Auffassung, welche die für endliche Reihen gültigen Rechnungsregeln sorglos auch auf die unendlichen Reihen übertrug: er setzte unmittelbar in Evidenz, daß die nicht absolut konvergenten Reihen bei gehöriger Gliederumstellung jede beliebige Summe ergeben können.

Besonderen Stolz legte Dirichlet auf seine Methode des diskontinuierlichen Faktors zur Bestimmung vielfacher Integrale. Er pflegte zu sagen, es ist das ein sehr einfacher Gedanke, und schmunzelnd hinzuzufügen, aber man muß ihn haben. Die Methode gestattet bekanntlich, sich über die Grenzen eines Integrationsraumes hinwegzusetzen durch Verwendung eines zweckmäßig geformten Faktors, der im Innern des Raumes 1 und außerhalb 0 ist. Wenn der Gedanke, aus 0 und 1, dem Nichts und dem All, nach dem dyadischen System die ganze Größenwelt aufzubauen, Leibniz schon derart gefiel, daß er sich dadurch eine Bekehrung des damaligen für Wissenschaft interessierten Kaisers von China vom Heidentume versprach, welche Hoffnungen auf die Vervollkommnung der Welt hätte nicht Leibniz aus einer Kenntnis des Dirichletschen diskontinuierlichen Faktors geschöpft.

Dirichlets Verdienste um die mathematische Physik sind besonders hoch darum einzuschätzen, daß er durch seine Vorlesungen außerordentlich für die Verbreitung dieser Lehren in Deutschland gewirkt hat. Gewaltige Offenbarungen auf diesem Gebiete hätten sich an seinen Namen geknüpft, wenn nicht sein Leben so jäh abgebrochen wäre. Er hatte einen mathematisch strengen Beweis für die Stabilität unseres Planetensystems gefunden, und er war in den Besitz einer ganz neuen, allgemeinen Methode der Behandlung und Auflösung der Differentialgleichungen der Mechanik gelangt. Aber es sind kaum Andeutungen darüber verblieben, denn er entschloß sich immer nur schwer zu der Niederschrift des Erforschten.

Gerade seine Göttinger Zeit war besonders ausgefüllt durch Nachdenken über physikalische Fragen. Und gerade im Hinblick auf die angewandten Gebiete der Mathematik hatte auch Wilhelm Weber Dirichlets Berufung hierher so warm befürwortet. Ich möchte eine allgemein gehaltene Stelle aus Webers Bericht wiedergeben, die hierzu als Einleitung diene und die Ihr Interesse erwecken wird:

Für die anregende Kraft, welche die Forschungen der einen Wissenschaft auf die der anderen ausübt, zeigt sich die enge Verwandtschaft

allein nicht entscheidend, es kommt vielmehr auf die Verschiedenartigkeit und gegenseitige Ergänzung der von zwei Seiten dargebotenen Elemente an. Der Astronom, welcher sich mit einem Physiker, der Physiker, welcher sich mit einem Astronomen zu einer physikalischen oder astronomischen Forschung verbinden wollte, würde keine wesentlich neuen Elemente dazu mitbringen. Der Mathematiker kann dagegen alle Reichtümer seiner Wissenschaft und die Resultate seiner eigenen Forschungen bei einer schicklich gewählten astronomischen oder physikalischen Untersuchung, zu der er sich mit einem Astronomen oder Physiker verbindet, entfalten und gewinnt dadurch selbst wieder Antrieb und Anregung zur Erforschung neuer Gebiete mathematischer Probleme. Gauß, so fährt Weber fort, hat keine Opfer gescheut, um sich selbst aller Elemente der praktischen Astronomie vollkommen zu bemächtigen; er hätte dasselbe in Beziehung auf die Physik vermocht: nur die Vermeidung von Störungen in seinen mathematischen Forschungen, wo der Verlust noch größer gewesen sein würde, sind der Grund seiner Verbindung mit mir zu physikalischen Untersuchungen gewesen, die unter seiner Leitung zu so großen Resultaten geführt haben. Nicht bloß zum Fortbau der höheren Mathematik also, sondern auch zu dem der Astronomie und Physik ist ein Mathematiker hervorragender Art das größte Bedürfnis.

In seinen Vorlesungen behandelte Dirichlet mit Vorliebe diejenigen Gebiete, an deren Ausbau er selbst reichen Anteil hatte. Sein Vortrag war dadurch so eindringlich, weil es schien, als wenn er im Begriffe stünde, eben erst den ganzen Bau zu schaffen; es war in hohem Grade fesselnd, ihm bei dieser Arbeit zu folgen. Er entwickelte den Stoff in vollster Natürlichkeit. Kein Kunstgriff trat auf als *deus ex machina*, tragisch geschürzte Knoten zu unerwartet günstiger Lösung zu führen.

So außerordentlich anregend Dirichlet als Lehrer gewirkt hat, eine besondere mathematische Schule hat er nicht gegründet. Aber viele, die sich hernach auf individuell sehr verschiedene Wege zerstreuten, verdanken ihm die stärksten Impulse ihres wissenschaftlichen Strebens. Seinen Enthusiasmus für die Anregungen Dirichlets meinte der junge Eisenstein nicht warm genug schildern zu können, auch wenn ihm ein ehernes Herz und eine tausendfältige Zunge verliehen wären. Welcher Mathematiker hätte kein Verständnis dafür, daß die leuchtende Bahn Riemanns, dieses riesigen Meteors am mathematischen Himmel, vom Sternbilde Dirichlets ihren Ausgangspunkt nahm. Mag auch das von Riemann Dirichletsches Prinzip benannte scharfe Schwert zuerst von William Thomsons jungem Arm geschwungen sein, von dem anderen Dirichletschen Prinzip, mit einem Minimum an blinder Rechnung, einem Maxi-

mum an sehenden Gedanken die Probleme zu zwingen, datiert die Neuzeit in der Geschichte der Mathematik. Niemals vergaß Kronecker zu sagen, wieviel von seiner mathematischen Existenz er Dirichlet schulde, wenn auch Dirichlet selbst Kronecker nur in die *unteren* Regionen *einer* der Wissenschaften eingeführt haben wollte, auf deren Höhen dieser als Meister einherschreite. Erst hier in Göttingen trat Dedekind in Beziehung zu Dirichlet; wir verehren in ihm den einzigen uns übriggebliebenen Heros aus der größten Epoche in der Arithmetik. Ganz in den Ideenkreis von Dirichlet trat auch Lipschitz ein, der in seinen jüngeren Jahren zu Helmholtz und später zu Hertz seine hohe Begeisterung für Dirichlet kundgab. Alle diese Männer von unvergleichlichen Verdiensten um die heutige Mathematik, den besten Teil ihrer mathematischen Kraft gewannen sie durch Dirichlet, und wir heute, wenn wir uns mehr als je bemühen, die Wissenschaft in ihrer einfachen Wahrheit zu erkennen und darzustellen, stehen wir nicht in der Schule Dirichlets?

Sachregister.

- Achse des Impulses eines Körpers in einer Flüssigkeit** II 288.
Adhäsionskraft II 338.
adjungierte Darstellungen quadratischer Formen I 91.
 — quadratische Formen I 80.
 — Genera I 83.
 — Gruppen von Darstellungen I 93.
 — Klassen quadratischer Formen I 82.
 — Substitutionen I 81.
äquivalente Darstellungen quadratischer Formen I 88.
 — Gruppen von Darstellungen I 90.
 — Parallelepipeda I 289.
Äquivalenz von linearen Formensystemen II 53.
 — von quadratischen Formen I 5, 210, 253, 321, II 21.
ähnliche Substitutionen I 205.
algebraische Zahl I 302, 357.
alternierende Matrix II 376.
Außenoberfläche II 123.
äußere Multiplikation I 251.
äußerer Raum einer geschlossenen Fläche II 123.
äußeres Volumen eines Körpers I 272.
äußerstes Paralleleipedum I 281.
Bedingung der Stützebene mit der äußeren Normalen (α , β , γ) II 233.
Beschleunigungsvektor II 439.
bewegender Kraftvektor II 401, 440.
Bewegungsvektor II 439.
Breite eines Körpers II 105, 149, 277.
Charaktere eines Genus quadratischer Formen I 7, 42, 74—75, 150, 161, 228.
charakteristische Form einer Klasse I 72, 161.
Darstellbarkeit, ganzzahlige, einer quadratischen Form durch eine zweite I 86.
 — — einer Zahl durch eine quadratische Form I 94.
 —, rationale, der Null durch quadratische Formen I 227.
Darstellung einer Zahl durch eine Summe von fünf Quadraten I 133.
Determinantenfläche II 73.
Diagonalkettenbruch I 344, II 46.
Dichtigkeit, mittlere, eines Punktgitters I 248, II 75.
Dirichletsche Reihen I 130, 189, II 82.
Diskriminante eines algebraischen Körpers I 256, 262, 276.
Diskontinuitätsbereich für die Äquivalenz quadratischer Formen II 54.
Distanzfunktion eines konvexen Körpers II 132.
duale Matrix (Raum-Zeit-Vektor) II 376.
Durchmesser eines Körpers II 287.
Eckpunkt eines konvexen Bezirks II 136, 163.
Eckstützebene II 165, 169, 259.
Eichbezirk II 221.
Eichfläche der Distanzen II 123.
Eichkörper der Strahldistanzen I 272.
eigentliche Äquivalenz quadratischer Formen I 210, 253.
 — Darstellung einer quadratischen Form durch eine zweite I 87.
 — — einer Zahl durch eine quadratische Formen I 94.
Eigenzeit eines Raum-Zeitfadens II 394.
Einheiten algebraischer Zahlkörper I 268, 276, 288, 312, 316, 362.
 — kubischer Zahlkörper II 47.
Einheitsmatrix II 376.
einhellige Strahldistanzen I 273.
Energiesatz II 401.
Enthalte sein einer quadratischen Form in einer zweiten I 9, 253.
extreme Form I 156, II 75.
 — — in bezug auf den reduzierten Raum II 76.
 — Formenklasse II 75.
 — Stützebene II 164, 233.
 — Parallelepipeda II 48.
 — Zylinder II 47.
extremer Halbraum um einen konvexen Körper II 233.
 — Punkt eines konvexen Bezirks II 157.
Fläche der Distanz II 123.
Flächendichte der Energie usw. II 346.
Flächenelement II 212.

- Flächeninhalt = Inhalt.
 Flächenpunkt eines konvexen Bezirks II 163.
 Formen mit größtem Minimum I 156.
 freies Parallelepipedium I 281.
 Fundamentalgleichung für die Teilungsfläche II 347.
 Fundamentalreihe für eine Kette I 289.

 Gaußsche Summen, mehrfache I 46, 229.
 Gegenpunkt II 4.
 zu Lösungen (b_h) gehörende Darstellungen quadratischer Formen I 97.
 zu ξ, η, ζ gehörende Substitution I 282.
 zu r gehörende Substitution I 303.
 zu einem Parallelepipedium gehörige Koordinaten II 7.
 gemischter Flächeninhalt II 194, 245.
 gemischtes Volumen dreier Körper II 129, 196, 241.
 Genus von quadratischen Formen I 7, 71, 149, 158, 221.
 Geometrie der Zahlen I 264, II 43.
 Geschlecht = Genus.
 geschlossene Fläche II 123.
 Gitter = Zahlengitter.
 gitterförmige Anordnung II 3.
 Gitteroktaeder II 9.
 — erster Art und zweiter Art II 11.
 Gitterpunkt I 281, 294, 321, II 120.
 gleichgestellte quadratische Formen I 146, II 57.
 Gleichzeitigkeit II 365.
 Grenze von Körpern II 235.
 Grenzformen I 155.
 Größe des Impulses eines Körpers in einer Flüssigkeit II 288.
 Grundform einer Klasse quadratischer Formen für einen Modul I 34, 183, 230.
 Grundgleichungen für die elektromagnetischen Vorgänge in bewegten Körpern II 367, 385.
 Grundparallelepiped eines Gitters II 5.
 Gruppen von Darstellungen einer quadratischen Form durch eine zweite I 90.
 — von Formen für einen Modul I 38.
 Gruppen linearer ganzzahliger Substitutionen I 208, 214.

 Halbraum um einen konvexen Bezirk II 163, 233.
 Hauptachse eines Körpers II 290.
 Hauptformenrest = Hauptrest einer Formenklasse I 25, 171, 231.
 Hauptideal I 257.

 Hauptrepräsentant einer Formenklasse I 25.
 eine Form ist höher als eine zweite I 146, II 57.
 homothetische Körper II 149, 245.
 hydrodynamischer Druck II 309.
 hydrostatischer Auftrieb II 317.
 — Druck II 343.

 Ideal I 257, II 455.
 Impuls der Bewegung eines Körpers in einer Flüssigkeit II 285.
 Impulsvektor II 441.
 Index = Trägheitsindex.
 Inhalt einer Figur I 325, 339.
 Innenoberfläche II 123.
 innerer Raum einer geschlossenen Fläche II 123.
 inneres Produkt zweier Strecken I 250.
 — Volumen eines Körpers I 272.
 Integral der mittleren Krümmung II 241.
 Invarianten σ, σ einer Ordnung von quadratischen Formen I 12.
 — eines Geschlechts von quadratischen Formen I 150, 160, 228.
 inwendige Bezirke einer Schar II 179.

 Kammern, äquivalente II 61.
 Kanten des reduzierten Raumes II 68.
 Kantenform II 69.
 Kantenpunkt eines konvexen Bezirks II 163.
 Kapillardruck II 310.
 kapillarer Auftrieb II 317.
 Kapillarität II 299.
 Kappe II 171.
 Kappenkörper II 175, 227, 259.
 Kette von extremen Zylindern II 47.
 — — Parallelepipeda I 284.
 — — Substitutionen I 303, 358.
 — zu Formen I 334.
 Kettenbrüche I 276, 278, 293, 320, II 44.
 kettenbruchähnliche Algorithmen I 243.
 Kettenglied I 334.
 kinetischer Druck II 309.
 kinetische Energie II 441.
 Klasse von quadratischen Formen I 10, 243, 254.
 Klassenanzahl quadratischer Formen II 95.
 Kohäsion II 299.
 Kohäsionskraft II 332.
 Kohäsionssprung II 309.
 Komponenten eines Raum-Zeit-Vektors I.
 Art II 377; II. Art II 378.
 Kongruenz zweier Formenklassen I 37, 137, 149, 158.
 konvexe Fläche II 123, 134.

- konvexer Bezirk II 135, 154.
 — Körper II 103, 123, 131, 232.
 — Restbereich II 120.
 — Zug II 110.
 konvexes Gebilde II 269.
 — Polyeder II 104, 233.
 Körper der Projektionen II 216.
 — — vierten Einheitswurzel II 49.
 — konstanter Breite II 277.
 — konstanten Umfanges II 278.
 Kraftvektor der Bewegung II 441.
 Kriterium für algebraische Zahlen I 302, II 50.
 Krümmungsfunktion eines Körpers II 263.
 Krümmungshyperbel II 440.
 Länge des Umfanges eines vollkommenen Ovals II 245.
 — einer geraden Linie I 272, II 44.
 — einer Kurve II 122.
 Leitungsstrom II 367, 371, 421.
 Lichtpunkt II 402.
 Linienpunkt eines konvexen Bezirks II 164.
 lor II 383, 409.
 Lorentz-Transformation II 360, 362.
 Maß der Darstellung einer Zahl durch quadratische Formen I 116.
 — einer Klasse positiver quadratischer Formen I 116, 158.
 — eines Genus I 116, 134, 150.
 — einer Ordnung I 116.
 — indefiniter Formen I 153.
 Massendichte II 395.
 Massenwirkung II 397.
 Matrix II 375.
 Minimum einer quadratischen Form I 154, 254, 270.
 mittlere Koeffizienten einer Form I 6.
 mittlere Dichtigkeit eines Punktsystems I 248, II 75.
 — Krümmung II 303.
 mittlerer Krümmungsradius II 209.
 — Stützebenenabstand II 209.
 Mittelwert II 212.
 ξ -Nachbar usw. eines Parallelepipeds I 283.
 eine Form ist niedriger als eine zweite I 146, II 57.
 Niveauebene II 305.
 Norm eines Ideals I 257.
 Normalensektor II 169, 172.
 normale Raum-Zeit-Vektoren II 378.
 Normalquerschnitt eines Raum-Zeitfadens II 394.
 Obere Normale eines Querschnitts II 394.
 Oberfläche einer krummen Fläche II 122, 256.
 — eines Polyeders II 206.
 Oberflächenenergie, molekulare II 350.
 Oberflächenentropie, molekulare II 350.
 Oberflächenintegral II 212.
 Oberflächenspannung II 299, 347.
 Okt ($\mathfrak{ABC}$) II 9.
 Ordnung der quadratischen Formen I 5, 12, 162, 228.
 — einer Gruppe von linearen Substitutionen I 209.
 Ortszeit II 437.
 Oval II 154, 242.
 —, vollkommenes II 242.
 Ovaloid, vollkommenes II 233.
 Parallelkettenbruch I 344.
 Periode eines Kettenbruches I 347.
 periodische Kette von Substitutionen I 358.
 periodischer Diagonalkettenbruch I 346.
 polarer Körper II 146.
 Polarisation, dielektrische II 423.
 Polyeder II 124, 136, 233.
 Polyederschar II 184.
 Polygon II 136.
 Postulat der absoluten Welt II 437.
 primäre Darstellung einer quadratischen Form I 102.
 primitive quadratische Formen I 12, 162.
 — — in bezug auf einen Modul I 12.
 — Grenzform I 155.
 Produkt von Idealen I 257.
 — von Matrizen II 375.
 Projektion eines Körpers auf eine Ebene II 215.
 Projektionsraum eines Körpers II 161.
 Projektionssektor eines Körpers II 161.
 Punktsystem, regelmäßiges I 248.
 Quadratisches System einer Form I 5.
 Querschnitt eines Raum-Zeitfadens II 394.
 Radius eines Körpers II 287.
 Randbezirke einer Schar II 179.
 Randwinkel II 306.
 Raumvektor Magnetisierung II 425.
 — Polarisation II 423.
 Raum-Zeitfaden II 394.
 — — linien II 393.
 — — -Matrix II 383.
 — — punkt II 355.
 — — -Sichel II 395.
 — — -Vektor I. und II. Art II 363, 364.
 — — — Geschwindigkeit II 369.

- reduzierte quadratische Formen I 146,
 154, 217, 244, 269, II 59.
 — — —, binäre II 23, 61.
 — — —, ternäre II 23.
 — — —, quaternäre I 146.
 — — —, quinäre I 154.
 — — —, senäre I 217.
 reduzierter Raum II 61.
 reduziertes Grundparallelepiped II 8.
 Reduktion des Zahlengitters in bezug auf
 ein Parallelepiped II 8.
 regelmäßiges Punktsystem I 248.
 Relativitätsprinzip, -postulat, -theorem
 II 353, 369.
 Relativoberfläche II 262.
 Repräsentant einer Formenklasse I 10.
 Rest einer Formenklasse I 13.
 Reste einer Form I 157.
 Restbereich II 120.
 reziproke Matrix II 376.
 Richtung II 133.
 ripples II 328.
 Röntgen-Vektor II 423.
 Ruh-Dichte der Elektrizität II 362, 382,
 409.
 Ruh-Kraft, elektrische II 380, 421.
 — —, magnetische II 381, 424.
 Ruh-Magnetisierung II 425.
 Ruh-Massendichte II 395.
 Ruh-Polarisation, dielektrische II 423.
 Ruh-Strahl II 382.
 Ruh-Strom II 382.
 auf Ruhe transformieren (einen Raum-
 Zeitpunkt) II 362, 393.
 Schar konvexer Bezirke II 179.
 — — Körper II 180, 239.
 Schwerpunkt II 143.
 Seite (α , β , γ) einer Ebene II 136.
 Seitenfläche eines Polyeders (ν -te) II 107.
 seitliche Koeffizienten einer Form I 6.
 singulärer Raum-Zeit-Vektor II 364.
 Spannungswirkung II 399.
 Steighöhe II 315.
 stetig gekrümmter konvexer Körper II 262.
 Strahldistanz I 272.
 Strahlgebilde eines Raum-Zeitpunktes
 II 402.
 Stützebene II 106, 136, 232, 277.
 Stützebenenabstand, mittlerer II 209, 213.
 Stützebenenfunktion II 4, 145, 232.
 Stützgerade II 154.
 Stützgeradenfunktion II 242.
 Substitution einer Kette I 334.
 Summe der Kantenkrümmung II 209.
 Tangentialebene an einen konvexen Be-
 zirk II 166.
 Tangentialkörper II 224.
 tangentialer Parameter II 107, 267.
 Teiler der Darstellung einer quadratischen
 Form durch eine zweite I 102.
 Teilungsfläche II 346.
 thermischer Druck II 343.
 transponierte Matrix II 375.
 Trägheitsindex I 5, 10, 160.
 Umfang eines Körpers II 278.
 — — Ovals II 207.
 unabhängige konvexe Körper II 129.
 — Richtungen II 6, 104.
 — Systeme ganzer Zahlen I 295, 302.
 Verallgemeinerte Oberfläche II 124.
 — Außen- und Innenoberfläche II 123.
 verallgemeinerter Flächeninhalt II 219.
 Verdampfungswärme II 340.
 Virial der Kohäsionskraft II 337.
 virtuelle Verrückung in der Raum-Zeit-
 Sichel II 396.
 vollkommene periodische Kette I 351.
 vollkommenes Oval II 242.
 — Ovaloid II 233.
 vollständige Äquivalenz von quadratischen
 Formen I 210.
 vollständiges Formensystem für ein Genus
 I 113.
 — Invariantensystem eines Genus I 150, 160.
 Volumen eines Körpers I 272, II 122, 143,
 241.
 — des reduzierten Raumes II 80.
 Wechselseitige Strahldistanzen I 273.
 Welt II 432.
 Weltlinie II 432.
 Weltpostulat II 437.
 Weltpunkt II 432.
 Winkeldistanz zweier Punkte II 209, 270.
 wirkliche Ecke I 251.
 Zahlen φ_k eines Hauptrepräsentanten I 33.
 Zahlengitter I 264, 271, 281, 294, 321,
 II 5, 43, 120.
 Zentrum der Hauptachsen eines Körpers
 II 286.
 zerfallende quadratische Formen I 223.
 Zerfallung einer Zahl in fünf Quadrate I 119.
 Zwischenhyperbeln II 438.

DIFFERENTIAL EQUATIONS

By H. BATEMAN

CHAPTER HEADINGS: I. Differential Equations and their Solutions. II. Integrating Factors. III. Transformations. IV. Geometrical Applications. V. Diff. Eqs. with Particular Solutions of a Specified Type. VI. Partial Diff. Eqs. VII. Total Diff. Eqs. VIII. Partial Diff. Eqs. of the Second Order. IX. Integration in Series. X. The Solution of Linear Diff. Eqs. by Means of Definite Integrals. XI. The Mechanical Integration of Diff. Eqs.

—1917-67. xi+306 pp. 5 $\frac{3}{8}$ x8.

[190] Prob. \$4.95

ERGODENTHEORIE, u.a.

By H. BEHNKE, P. THULLEN, and E. HOPF

TWO VOLUMES IN ONE.

THEORIE DER FUNKTIONEN MEHRERER KOMPLEXER VERAENDERLICHEN, by H. Behnke and P. Thullen.

ERGODENTHEORIE, by E. Hopf.

—1934-67; 1937-67. 210 pp. 5 $\frac{1}{2}$ x8 $\frac{1}{2}$.

[68A] Two vols. in one. In prep.

L'APPROXIMATION

By S. BERNSTEIN and CH. de LA VALLÉE POUSSIN

TWO VOLUMES IN ONE:

Leçons sur les Propriétés Extrémales et la Meilleure Approximation des Fonctions Analytiques d'une Variable Réelle, by Bernstein.

Leçons sur l'approximation des Fonctions d'une Variable Réelle, by Vallée Poussin.

—1925-66; 1919-66. vii + 204 + 152 pp. 6x9. [198]

Two vols. in one. In prep.

CONFORMAL MAPPING

By L. BIEBERBACH

"The first book in English to give an elementary, readable account of the Riemann Mapping Theorem and the distortion theorems and uniformisation problem with which it is connected. . . . Presented in very attractive and readable form."

—*Math. Gazette.*

"Engineers will profitably use this book for its accurate exposition."—*Appl. Mechanics Reviews.*

—1952. vi+234 pp. 4 $\frac{1}{2}$ x6 $\frac{1}{2}$.

[90] Cloth \$2.75
[176] Paper \$1.50

BASIC GEOMETRY

By G. D. BIRKHOFF and R. BEATLEY

A highly recommended high-school text by two eminent scholars.

"is in accord with the present approach to plane geometry. It offers a sound mathematical development . . . and at the same time enables the student to move rapidly into the heart of geometry."—*The Mathematics Teacher.*

"should be required reading for every teacher of Geometry."—*Mathematical Gazette.*

—Third edition. 1959. 294 pp. 5 $\frac{1}{4}$ x8.

[120] \$3.95

VORLESUNGEN UEBER DIFFERENTIALGEOMETRIE, Vols. I, II

By W. BLASCHKE

TWO VOLUMES IN ONE.

Partial Contents: VOL. I. 1. Theory of Curves. 2. Extremal Curves. 3. Strips. 4. Theory of Surfaces. 5. Invariant Derivatives on a Surface. 6. Geometry on a Surface. 7. On the Theory of Surfaces in the Large. 8. Extremal Surfaces. 9. Line Geometry.

VOL. II. *Affine Differential Geometry.* 1. Plane Curves in the Small. 2. Plane Curves in the Large. 3. Space Curves. 5. Theory of Surfaces. 6. Extremal Surfaces. 7. Special Surfaces.

—1930; 1923-67. x+311+ix+259 pp. 6x9.

[202] Two vols. in one. **Prob. \$9.50**

KREIS UND KUGEL

By W. BLASCHKE

Isoperimetric properties of the circle and sphere, the (Brunn-Minkowski) theory of convex bodies, and differential-geometric properties (in the large) of convex bodies. A standard work.

—x + 169 pp. 5½x8½.

[59] **\$3.50**

INTEGRALGEOMETRIE

By W. BLASCHKE and E. KÄHLER

THREE VOLUMES IN ONE.

VORLESUNGEN UEBER INTEGRALGEOMETRIE, Vols. I and II, by W. Blaschke.

EINFUEHRUNG IN DIE THEORIE DER SYSTEME VON DIFFERENTIALGLEICHUNGEN, by E. Kähler.

—222 pp. 5½x8½.

[64] Three Vols. in One **\$4.95**

VORLESUNGEN UEBER FOURIERSCH INTEGRALE

By S. BOCHNER

"A readable account of those parts of the subject useful for applications to problems of mathematical physics or pure analysis."

—*Bulletin of the A. M. S.*

—1932. 237 pp. 5½x8½. Orig. pub. at \$6.40. [42] **\$3.50**

ALMOST PERIODIC FUNCTIONS

By H. BOHR

Translated by H. COHN. From the famous series *Ergebnisse der Mathematik und ihrer Grenzgebiete*, a beautiful exposition of the theory of Almost Periodic Functions written by the creator of that theory.

—1951-66. 120 pp. Lithotyped.

[27] **\$2.75**

WISSENSCHAFTLICHE ABHANDLUNGEN

By L. BOLTZMANN

—1909. Repr. of 1st ed. $6\frac{1}{2} \times 9\frac{1}{4}$. Three vol. set. In prep.

LECTURES ON THE
CALCULUS OF VARIATIONS

By O. BOLZA

A standard text by a major contributor to the theory. Suitable for individual study by anyone with a good background in the Calculus and the elements of Real Variables. The present, second edition differs from the first primarily by the inclusion within the text itself of various addenda to the first edition, as well as some notational improvements.

—2nd (corr.) ed. 1961. 280 pp. $5\frac{3}{8} \times 8$. [145] Cloth \$3.25
[152] Paper \$1.19

VORLESUNGEN UEBER
VARIATIONSRECHNUNG

By O. BOLZA

A standard text and reference work, by one of the major contributors to the theory.

—1909-63. C. repr. of 1st ed. ix+715 pp. $5\frac{3}{8} \times 8$. [160] \$8.00

THEORIE DER KONVEXEN KOERPER

By T. BONNESEN and W. FENCHEL

"Remarkable monograph."

—J. D. Tamarkin, *Bulletin of the A. M. S.*

—1934. 171 pp. $5\frac{1}{2} \times 8\frac{1}{2}$. Orig. publ. at \$7.50 [54] \$3.95

THE CALCULUS OF FINITE DIFFERENCES

By G. BOOLE

A standard work on the subject of finite differences and difference equations by one of the seminal minds in the field of finite mathematics.

Numerous exercises with answers.

—Fourth edition. 1958. xii + 336 pp. 5×8 . [121] Cloth \$3.95
[148] Paper \$1.39

A TREATISE ON
DIFFERENTIAL EQUATIONS

By G. BOOLE

Including the Supplementary Volume.

—Fifth edition. 1959. xxiv + 735 pp. $5\frac{1}{4} \times 8$. [128] \$6.00

CAJORI, "History of Slide Rule," see Ball

INTRODUCTORY TREATISE ON LIE'S THEORY OF FINITE CONTINUOUS TRANSFORMATION GROUPS

By J. E. CAMPBELL

Partial Contents: CHAP. I. Definitions and Simple Examples of Groups. II. Elementary Illustrations of Principle of Extended Point Transformations. III. Generation of Group from Its Infinitesimal Transformations. V. Structure Constants. VI. Complete Systems of Differential Equations. VII. Diff. Eqs. Admitting Known Transf. Groups. VIII. Invariant Theory of Groups. IX. Primitive and Stationary Groups. XI. Isomorphism. XIII. Construction of Groups from Their Structure Constants . . . XIV. Pfaff's Equation . . . XV. Complete Systems of Homogeneous Functions. XVI. Contact Transformations. XVII. Geometry of C. T.'s. XVIII. Infinitesimal C. T.'s. XX. Differential Invariants. XXI-XXIV. Groups of Line, of Plane, of Space. XXV. Certain Linear Groups.

—1903-66. xx + 416 pp. 5 $\frac{3}{8}$ x8.

[183] \$6.50

THEORY OF FUNCTIONS

By C. CARATHÉODORY

Translated by F. STEINHARDT. The recent, and already famous textbook, *Funktionentheorie*.

Partial Contents: **Part One.** Chap. I. Algebra of Complex Numbers. II. Geometry of Complex Numbers. III. Euclidean, Spherical, and Non-Euclidean Geometry. **Part Two.** Theorems from Point Set Theory and Topology. Chap. I. Sequences and Continuous Complex Functions. II. Curves and Regions. III. Line Integrals. **Part Three.** Analytic Functions. Chap. I. Foundations. II. The Maximum-modulus principle. III. Poisson Integral and Harmonic Functions. IV. Meromorphic Functions. **Part Four.** Generation of Analytic Functions by Limiting Processes. Chap. I. Uniform Convergence. II. Normal Families of Meromorphic Functions. III. Power Series. IV. Partial Fraction Decomposition and the Calculus of Residues. **Part Five.** Special Functions. Chap. I. The Exponential Function and the Trigonometric Functions. II. Logarithmic Function. III. Bernoulli Numbers and the Gamma Function.

Vol. II.: Part Six. Foundations of Geometric Function Theory. Chap. I. Bounded Functions. II. Conformal Mapping. III. The Mapping of the Boundary. **Part Seven.** The Triangle Function and Picard's Theorem. Chap. I. Functions of Several Complex Variables. II. Conformal Mapping of Circular-Arc Triangles. III. The Schwarz Triangle Functions and the Modular Function. IV. Essential Singularities and Picard's Theorems.

"A book by a master . . . Carathéodory himself regarded [it] as his finest achievement . . . written from a catholic point of view."—*Bulletin of A.M.S.*

—Vol. I. Second edition, 1958. 310 pp. 6x9.

[97] \$4.95

—Vol. II. Second edition, 1960. 220 pp. 6x9.

[106] \$4.95

ALGEBRAIC THEORY OF MEASURE AND INTEGRATION

By C. CARATHÉODORY

Translated from the German by FRED E. J. LINTON. By generalizing the concept of point function to that of a function over a Boolean ring ("soma" function), Prof. Carathéodory gives an algebraic treatment of measure and integration.

CONTENTS: CHAP. I. Somas (Axiomatic method, somas as elements of a Boolean ring, . . .). II. Set of Somas. III. Place Functions (Functionoids). IV. Computation with Place Functions. V. Measure Functions. VI. The Integral. VII. Application of Integration to Limit Processes. VIII. Computation of Measure Functions. IX. Regular Measure Functions. X. Isotypic Regular Measure Functions. XI. Content Functions. APPENDIX: Somas as Elements of Partially Ordered Sets.

—1963. 378 pp. 6x9.

[161] \$7.50

VORLESUNGEN UBER REELLE FUNKTIONEN

By C. CARATHÉODORY

This great classic is at once a book for the beginner, a reference work for the advanced scholar and a source of inspiration for the research worker.

—In prep. 2nd ed. 5 $\frac{3}{8}$ x8. 728 pp.

[38] Prob. \$9.50

CARSLAW, "Non-Euclidean Plane Geometry," see Ball

ELECTRIC CIRCUIT THEORY and the OPERATIONAL CALCULUS

By J. R. CARSON

"A rigorous and logical exposition and treatment of the Heaviside operational calculus and its applications to electrical problems . . . will be enjoyed and studied by mathematicians, engineers, and scientists."—*Electrical World*.

—2nd ed. 206 pp. 5 $\frac{3}{8}$ x8.

[92] \$3.95

COLLECTED PAPERS (OEUVRES)

By P. L. CHEBYSHEV

One of Russia's greatest mathematicians, Chebyshev (Tchebycheff) did work of the highest importance in the Theory of Probability, Number Theory, and other subjects. The present work contains his post-doctoral papers (sixty in number) and miscellaneous writings. The language is French, in which most of his work was originally published; those papers originally published in Russian are here presented in French translation.

—1962. Repr. of 1st ed. 1,480 pp. 5 $\frac{3}{8}$ x8 $\frac{1}{4}$.

[157] Two Vol. set. \$27.50

THEORIE GENERALE DES SURFACES

By G. DARBOUX

One of the great works of the mathematical literature.

An unabridged reprint of the latest edition of *Leçons sur la Théorie générale des surfaces et les applications géométriques du Calcul infinitésimal*.

—Vol. I (2nd ed.) xii+630 pp. Vol. II (2nd ed.) xvii+584 pp.
Vol. III (1st ed.) xvi+518 pp. Vol. IV. (1st ed.) xvi+537 pp.
[216] In prep.

ALGEBREN

By M. DEURING and W. KRULL

TWO VOLUMES IN ONE.

ALGEBREN, by *Deuring*.

IDEALTHEORIE, by *Krull*.

—1935-68. v+143+152 pp. 5½x8½. [50A] In prep.

COLLECTED PAPERS

By L. E. DICKSON

—1968. 4 vols. Approx. 3,400 pp. 6½x9¼. In prep.

HISTORY OF THE THEORY OF NUMBERS

By L. E. DICKSON

"A monumental work . . . Dickson always has in mind the needs of the investigator . . . The author has [often] expressed in a nut-shell the main results of a long and involved paper in a much clearer way than the writer of the article did himself. The ability to reduce complicated mathematical arguments to simple and elementary terms is highly developed in Dickson."—*Bulletin of A. M. S.*

—Vol. I (Divisibility and Primality) xii+486 pp. Vol. II (Diophantine Analysis) xxv+803 pp. Vol. III (Quadratic and Higher Forms) v+313 pp. [86] Three vol. set \$19.95

STUDIES IN THE THEORY OF NUMBERS

By L. E. DICKSON

A systematic exposition, starting from first principles, of the arithmetic of quadratic forms, chiefly (but not entirely) ternary forms, including numerous original investigations and correct proofs of a number of classical results that have been stated or proved erroneously in the literature.

—1930-62 viii+230 pp. 5¾x8. [151] \$3.95

ALGEBRAIC NUMBERS

By L. E. DICKSON, et al.

TWO VOLUMES IN ONE.

Both volumes of the *Report of the Committee on Algebraic Numbers* are here included, the authors being L. E. Dickson, R. Fueter, H. H. Mitchell, H. S. Vandiver, and G. E. Wahlen.

Partial Contents: CHAP. I. Algebraic Numbers. II. Cyclotomy. III. Hensel's p -adic Numbers. IV. Fields of Functions. I'. The Class Number in the Algebraic Number Field. II'. Irregular Cyclotomic Fields and Fermat's Last Theorem.

—1923-67, 1928-67. 96+111 pp. 6x9. In prep.

THE INTEGRAL CALCULUS

By J. W. EDWARDS

A leisurely, immensely detailed, textbook of over 1,900 pages, rich in illustrative examples and manipulative techniques and containing much interesting material that must of necessity be omitted from less comprehensive works.

There are forty large chapters in all. The earlier cover a leisurely and a more-than-usually-detailed treatment of all the elementary standard topics. Later chapters include: Jacobian Elliptic Functions, Weierstrassian Elliptic Functions, Evaluation of Definite Integrals, Harmonic Analysis, Calculus of Variations, etc. Every chapter contains many exercises (with solutions).

—2 vols. 1,922 pp. 5x8. [102], [105] Each volume \$8.50

TRANSFORMATIONS OF SURFACES

By L. P. EISENHART

Many of the advances in the differential geometry of surfaces in the present century have had to do with transformations of surfaces of a given type into surfaces of the same type. The present book studies two types of transformation to which many, if not all, such transformations can be reduced.

—2nd (Corr.) ed. ix+379 pp. 5 $\frac{3}{8}$ x8. [167] \$4.95

MATHEMATISCHE ABHANDLUNGEN

By G. EISENSTEIN

—Repr. of edition of 1847. iv+336 pp. 6 $\frac{1}{2}$ x9 $\frac{1}{4}$. In prep.

INTRODUCTION TO THE ALGEBRA OF QUANTICS

By E. B. ELLIOTT

—Repr. of 2nd ed. 1913-64 xvi+416 pp. 5 $\frac{3}{8}$ x8 [184] \$6.00

ASYMPTOTIC SERIES

By W. B. FORD

TWO VOLUMES IN ONE: *Studies on Divergent Series and Summability* and *The Asymptotic Developments of Functions Defined by MacLaurin Series*.

PARTIAL CONTENTS: I. MacLaurin Sum-Formula; Introduction to Study of Asymptotic Series. II. Determination of Asymptotic Development of a Given Function. III. Asymptotic Solutions of Linear Differential Equations. . . . V. Summability, etc. I. First General Theorem. . . . III. MacLaurin Series whose General Coefficient is Algebraic. . . . VII. Functions of Bessel Type. VIII. Asymptotic Behavior of Solution of Differential Equations of Fuchsian Type. Bibliography.

—1916; 1936-60. x+341 pp. 6x9. [143] Two vols. in one. \$6.00

AUTOMORPHIC FUNCTIONS

By L. R. FORD

"Comprehensive . . . remarkably clear and explicit."—*Bulletin of the A. M. S.*

—2nd ed. (Cor. repr.) x+333 pp. 5 $\frac{3}{8}$ x8. [85] \$6.00

GYROSCOPIC THEORY

By G. GREENHILL

This work is intended to serve as a collection in one place of the various methods of the theoretical explanation of the motion of a spinning body, and as a reference for mathematical formulas required in practical work.

Originally published as a report to the British Advisory Committee for Aeronautics.

CHAPTER HEADINGS: I. Steady Gyroscopic Motion. II. Gyroscopic Applications. III. General Unsteady Motion of a Top or Gyroscope. IV. Geometrical Representation of the Motion of a Top. V. Algebraical Cases of Top Motion. VI. Numerical Illustrations and Diagrams. VII. The Spherical Pendulum. VIII. Motion referred to Moving Origin and Axes. IX. Dynamical Problems of Steady Motion and Small Oscillation. *Fold-out plates.*

—1914-67. vi + 277 + Plates. $6\frac{1}{2} \times 10\frac{3}{4}$.

[205] \$9.50

COMMUTATIVE NORMED RINGS

By I. M. GELFAND, D. A. RAIKOV, and G. E. SHILOV

Translated from the Russian. In the second English edition—to appear in 1967—Chapter II has been revised in accordance with a manuscript especially prepared for this edition by Professor Shilov.

Partial Contents: CHAPS. I AND II. General Theory of Commutative Normed Rings. III. Ring of Absolutely Integrable Functions and their Discrete Analogues. IV. Harmonic Analysis on Commutative Locally Compact Groups. V. Ring of Functions of Bounded Variation on a Line. VI. Regular Rings. VII. Rings with Uniform Convergence. VIII. Normed Rings with an Involution and their Representations. IX. Decomposition of Normed Ring into Direct Sum of Ideals. HISTORICO-BIBLIOGRAPHICAL NOTES. BIBLIOGRAPHY.

—2nd ed. 1967. 306 pp. 6x9.

[170] \$6.50

LES INTEGRALES DE STIELTJES et leurs Applications aux Problèmes de la Physique Mathématique

By N. GUNTHER

The present work is a reprint of Vol. I of the publications of the V. A. Steklov Institute of Mathematics, in Moscow. The text is in French.

—1932-49. 498 pp. $5\frac{3}{8} \times 8$.

[63] \$6.50

ONDES: Leçons sur la Propagation des Ondes et les Equations de l'Hydrodynamique

By J. HADAMARD

"[Hadamard's] unusual analytic proficiency enables him to connect in a wonderful manner the physical problem of propagation of waves and the mathematical problem of Cauchy concerning the characteristics of partial differential equations of the second order."—*Bulletin of the A. M. S.*

—viii + 375 pp. $5\frac{1}{2} \times 8\frac{1}{2}$.

[58] \$4.95

GRUNDZUEGE DER MENGENLEHRE

By F. HAUSDORFF

Some of the topics in the Grundzüge omitted from later editions:

Symmetric Sets—Principle of Duality—most of the "Algebra" of Sets—most of the "Ordered Sets"—Partially Ordered Sets—Arbitrary Sets of Complexes—Normal Types—Initial and Final Ordering—Complexes of Real Numbers—General Topological Spaces—Euclidean Spaces—the Special Methods Applicable in the Euclidean plane—Jordan's separation Theorem—The Theory of Content and Measure—The Theory of the Lebesgue Integral.

—Repr. of original (1914) edition, 484 pp. $5\frac{1}{4} \times 8$. [61] \$6.00

SET THEORY

By F. HAUSDORFF

Hausdorff's classic text-book is an inspiration and a delight. The translation is from the Third (latest) German edition.

"We wish to state without qualification that this is an indispensable book for all those interested in the theory of sets and the allied branches of real variable theory."—*Bulletin of A. M. S.*

—2nd ed. 1962. 352 pp. 6x9.

[119] \$6.50

VORLESUNGEN UEBER DIE THEORIE DER ALGEBRAISCHEN ZAHLEN

By E. HECKE

"An elegant and comprehensive account of the modern theory of algebraic numbers."

—*Bulletin of the A. M. S.*

—1923. 264 pp. $5\frac{1}{2} \times 8\frac{1}{2}$.

[46] \$4.95

INTEGRALGLEICHUNGEN UND GLEICHUNGEN MIT UNENDLICHVIELEN UNBEKANNTEN

By E. HELLINGER and O. TOEPLITZ

"Indispensable to anybody who desires to penetrate deeply into this subject."—*Bulletin of A.M.S.*

—With a preface by E. Hilb. 1928. 286 pp. $5\frac{1}{4} \times 8$. [89] \$4.50

THEORIE DER ALGEBRAISCHE FUNKTIONEN EINER VARIABLEN

By K. HENSEL and G. LANDSBERG

Partial Contents: PART ONE (Chaps. 1-8): Algebraic Functions on a Riemann Surface. PART TWO (Chaps. 9-13): The Field of Algebraic Functions. PART THREE (Chaps. 14-22): Algebraic Divisors and the Riemann-Roch Theorem. PART FOUR (Chaps. 23-27): Algebraic Curves. PART FIVE (Chaps. 28-31): The Classes of Algebraic Curves. PART SIX (Chaps. 32-37): Algebraic Relations among Abelian Integrals. APPENDIX: Historical Development. Geometrical Methods. Arithmetical Methods.

—1902-55. xvi + 707 pp. 6x9.

[179] \$9.50

Grundzüge Einer Allgemeinen Theorie der LINEAREN INTEGRALGLEICHUNGEN

By D. HILBERT

—306 pp. $5\frac{1}{2} \times 8\frac{1}{4}$.

[91] \$4.50

GEOMETRY AND THE IMAGINATION

By D. HILBERT and S. COHN-VOSSEN

Translated from the German by P. NEMENYI.

"A fascinating tour of the 20th century mathematical zoo. . . . Anyone who would like to see proof of the fact that a sphere with a hole can always be bent (no matter how small the hole), learn the theorems about Klein's bottle—a bottle with no edges, no inside, and no outside—and meet other strange creatures of modern geometry will be delighted with Hilbert and Cohn-Vossen's book."

—*Scientific American*.

"Should provide stimulus and inspiration to every student and teacher of geometry."—*Nature*.

"A mathematical classic. . . . The purpose is to make the reader *see* and *feel* the proofs. . . . readers can penetrate into higher mathematics with . . . pleasure instead of the usual laborious study."

—*American Scientist*.

"Students, particularly, would benefit very much by reading this book . . . they will experience the sensation of being taken into the friendly confidence of a great mathematician and being shown the real significance of things."—*Science Progress*.

"A person with a minimum of formal training can follow the reasoning. . . . an important [book]."

—*The Mathematics Teacher*.

"A remarkable book. . . . A veritable geometric anthology. . . . Over 330 diagrams and every [one] tells a story."—*The Mathematical Gazette*.

—1952. 358 pp. 6x9.

[87] \$6.00

COLLECTED PAPERS

(Gesammelte Abhandlungen)

By D. HILBERT

The present work contains all of David Hilbert's papers on Mathematics, Physics, and the Foundations of Mathematics. An exception is his work on Integral Equations, which is available in a separate volume [*Integralgleichungen*, in German] and those of his papers on the Foundations of Mathematics that appear as an appendix in his book on the Foundations of Geometry [*Grundlagen der Geometrie*, in German; English translation in prep.].

Volume I (Number Theory) contains Hilbert's papers on Number Theory, including his long paper on Algebraic Numbers. Volume II (Algebra, Invariant Theory, Geometry) covers not only the topics indicated in the sub-title but also papers on Diophantine Equations. Volume III carries the sub-title: Analysis, Foundation of Mathematics, Physics, and Miscellaneous Papers.

—1966. Repr. of 1st ed. 1,457 pp. 6x9. [195] Each vol. \$9.50

Three vol. set. \$25.00

